

Contingency Games for Multi-Agent Interaction

Lasse Peters^{ID}, *Graduate Student Member, IEEE*, Andrea Bajcsy^{ID},
Chih-Yuan Chiu^{ID}, *Graduate Student Member, IEEE*, David Fridovich-Keil^{ID}, *Member, IEEE*, Forrest Laine^{ID},
Laura Ferranti^{ID}, and Javier Alonso-Mora^{ID}, *Senior Member, IEEE*

Abstract—Contingency planning, wherein an agent generates a set of possible plans conditioned on the outcome of an uncertain event, is an increasingly popular way for robots to act under uncertainty. In this work we take a game-theoretic perspective on contingency planning, tailored to multi-agent scenarios in which a robot's actions impact the decisions of other agents and vice versa. The resulting *contingency game* allows the robot to efficiently interact with other agents by generating strategic motion plans conditioned on multiple possible intents for other actors in the scene. Contingency games are parameterized via a scalar variable which represents a future time when intent uncertainty will be resolved. By estimating this parameter online, we construct a game-theoretic motion planner that adapts to changing beliefs while anticipating future certainty. We show that existing variants of game-theoretic planning under uncertainty are readily obtained as special cases of contingency games. Through a series of simulated autonomous driving scenarios, we demonstrate that contingency games close the gap between certainty-equivalent games that commit to a single hypothesis and non-contingent multi-hypothesis games that do not account for future uncertainty reduction.

Index Terms—Planning under uncertainty, human-aware motion planning, motion and path planning.

I. INTRODUCTION

IMAGINE you are driving and you see a pedestrian in the middle of the road as shown in Fig. 1. The pedestrian is likely to continue walking to the right, but you also saw them turning their head around; so maybe they want to walk back to the left? You think to yourself, If the pedestrian continues to the right, I just need to decelerate slightly and can safely pass on the

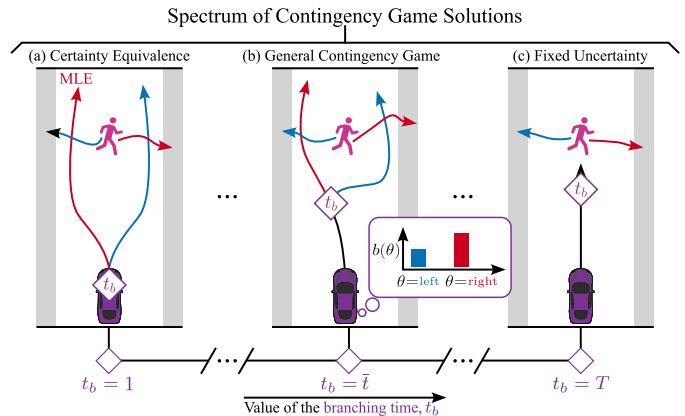


Fig. 1. Vehicle approaching a jaywalking pedestrian with uncertain intent. (a) A risk-taking driver may gamble for the most likely outcome, ignore uncertainty, and pass on the left. (c) A risk-averse driver may hedge against all outcomes by bringing their vehicle to a full halt, waiting for the situation to resolve. (b) An experienced driver may realize that this uncertainty will resolve in the near future (at time t_b) and thus commit to an *immediate plan* that can be continued safely and efficiently under both outcomes. Our contingency games formalize this middle ground between the two extremes.

left; but if they suddenly turn around, I need to brake and pass them on the right. Moreover, you understand that your actions influence the pedestrian's decision regarding whether and how quickly to cross the street. You decide to take your foot off the gas pedal and drive forward, aiming to pass the pedestrian on the left, but you are ready to brake and swerve to the right should the pedestrian turn around.

This example captures three important aspects of real-world multi-agent reasoning: (i) strategic interdependence of agents' actions due to their (partially) conflicting intents—e.g., the pedestrian's actions do not only depend on their own intent but also on your actions, and vice versa; (ii) accounting for uncertainty—e.g., how likely is it that the pedestrian wants to move left or right?; and (iii) planning contingencies—e.g., by anticipating that uncertainty will be resolved in the future, a driver can commit to an *immediate plan* (shown in black in Fig. 1(b)) that can be continued safely and efficiently under each outcome (shown in red and blue in Fig. 1(b)).

In this letter, we formalize this kind of reasoning by introducing *contingency games*: a mathematical model for planning contingencies through the lens of dynamic game theory. Specifically, we focus on model-predictive game-play (MPGP) settings, wherein an ego agent (i.e., a robot) plans by choosing its future trajectory in strategic *equilibrium* with those of other nearby agents (e.g., humans) at each planning invocation. Importantly, the ego agent must account for any uncertainty about other agents' intents—i.e., their optimization objectives—when

Manuscript received 17 August 2023; accepted 20 December 2023. Date of publication 16 January 2024; date of current version 26 January 2024. This letter was recommended for publication by Associate Editor Y. Hu and Editor G. Venture upon evaluation of the reviewers' comments. This work was supported in part by the European Union under the ERC Grant INTERACT 101041863, in part by the Dutch Science Foundation NWO-TTW within Veni project HARMONIA 18165, and in part by the National Science Foundation under Grant 2211548. (Corresponding author: Lasse Peters.)

Lasse Peters, Laura Ferranti, and Javier Alonso-Mora are with the Department of Cognitive Robotics, TU Delft, 2628 CD Delft, The Netherlands (e-mail: l.peters@tudelft.nl; l.ferranti@tudelft.nl; j.alonsomora@tudelft.nl).

Andrea Bajcsy is with Robotics Institute, Carnegie Mellon University, Pittsburgh, PA 15213 USA (e-mail: abajcsy@cmu.edu).

Chih-Yuan Chiu is with the Department of EECS, UC Berkeley, Berkeley, CA 94720 USA (e-mail: chihyuan_chiu@berkeley.edu).

David Fridovich-Keil is with the Department of Aerospace Engineering and Engineering Mechanics, UT Austin, Austin, TX 78712 USA (e-mail: dfk@utexas.edu).

Forrest Laine is with the Department of Computer Science, Vanderbilt University, Nashville, TN 37212 USA (e-mail: forrest.laine@vanderbilt.edu).

Website/Code: <https://lasse-peters.net/pub/contingency-games>

This letter has supplementary downloadable material available at <https://doi.org/10.1109/LRA.2024.3354548>, provided by the authors.

Digital Object Identifier 10.1109/LRA.2024.3354548

solving this game. For computational tractability, most established methods take one of two approaches. One class of methods ignores uncertainty by performing maximum likelihood estimation (MLE) and solving a certainty-equivalent game [1], [2], [3]. The other class accounts for uncertainty by planning with a full distribution—or “belief”—conservatively assuming that intent uncertainty will never be resolved during the planning horizon [4], [5]. Our **main contribution** is a game-theoretic interaction model that bridges the gap between these two extremes:

A contingency game is a model for *strategic interactions* which allows a robot to consider the full *distribution* of other agents’ intents while *anticipating intent certainty* in the near future.

Importantly, unlike existing formulations of trajectory games with parametric uncertainty [4], [5], [6], contingency games capture the fact that future belief updates will reduce uncertainty and hence, eventually, the *true* intent of the human will be clear at a future “branching” time, t_b . As a result, solutions of contingency games are *conditional* robot strategies which take a tree structure as shown in Fig. 1(b). The trunk of the conditional plan encodes decisions that are made before certainty is reached (before t_b). After t_b , the robot generates separate conditional trajectories for each possibility $\theta \in \Theta$.

Beyond our main contribution of an uncertainty-aware game-theoretic interaction model, we also (i) show how general-sum N -player contingency games can be transformed into mixed complementarity problems, for which off-the-shelf solvers [7] are available and (ii) discuss how beliefs and branching times may be estimated online for receding-horizon operation. We also highlight the desirable modeling flexibility of contingency games as a function of the parameter t_b , recovering certainty-equivalent games on one extreme, and conservative solutions on the other (see Fig. 1). Through a series of simulation experiments, we demonstrate that contingency games close the gap between these two extremes, and highlight the utility of estimating the branching time online.

II. RELATED WORK

A. Game-Theoretic Motion Planning

Game-theoretic planning has become increasingly popular in interactive robotics domains like autonomous driving [1], [8], [9], drone racing [10], and shared control [11] due to its ability to model influence among agents. A crucial axis in which prior works differ is in the modeling of the robot’s uncertainty regarding other agents’ objectives, dynamics, and state at each time. Methods which assume no uncertainty in the trajectory game model (e.g., taking the most probable hypothesis as truth) result in the simplest game formulations, and have been explored extensively [2], [12], [13], [14]. However, in real-world settings it is unrealistic to assume that a robot has full certainty, especially with respect to other agents’ intents. Instead, robots will often maintain a probability distribution, or “belief,” over uncertain aspects of the game.

There are several ways in which such a belief can be incorporated in a trajectory game. On one hand, the robot could simply optimize for its expected cost under the distribution of uncertain game elements. We call this a “fixed uncertainty” approach, since the game ignores the fact that as the game evolves, the robot could gain information leading to belief updates [4], [5], [6]. While these methods do utilize the robot’s uncertainty, they

often lead to overly conservative plans because the robot cannot reason about future information that would make it more certain (i.e., confident) in its decisions.

On the other hand, the agents in a game may reason about how their actions could lead to information gain; we refer to this as “dynamic uncertainty.” Games which exactly model dynamic information gain are inherently more complex, and are generally intractable to solve, especially when the belief space is large and has non-trivial update dynamics [15], [16]. Recent methods attempted to alleviate the computational burden of an exact solution via linear-quadratic-Gaussian approximations of the dynamics, objectives, and beliefs [17]. While the computational benefits of such approximations are significant, they introduce artifacts that make them inappropriate for scenarios such as that shown in Fig. 1 in which uncertainty is fundamentally multimodal. It remains an open challenge to solve “dynamic uncertainty” games tractably.

Our *contingency games* approach presents a middle ground between these paradigms via a belief-update model simple enough to compute exact solutions to the resulting dynamic-uncertainty game, and realistic enough to generate intelligent behavior that anticipates future uncertainty reduction.

B. Non-Game-Theoretic Contingency Planning

There is a growing literature of non-game-theoretic interaction planners [18], including various flavors of contingency planning. Online-optimization-based approaches primarily focus on predict-then-plan contingency planning [19], [20], [21], [22]. Recent learning-based contingency planners leverage deep neural networks to generate human predictions conditioned on candidate robot plans [23], [24]; a robot plan is then selected via methods like sampling-based model-predictive control [23], dynamic programming [25], or neural trajectory decoders [24].

Other approaches draw inspiration from deep reinforcement learning to accomplish both intention prediction and motion planning. To predict multi-agent trajectories in the near future, Packer et al. [26] and Rhinehart et al. [27] construct a flow-based generative model and a likelihood-based generative model, respectively. Meanwhile, Rhinehart et al. [28] develop an end-to-end contingency planning framework for both intention prediction and motion planning. We bring the notion of contingency planning to a different modeling domain—dynamic game theory—to extend its capabilities to handle “dynamic uncertainty” in strategic interactions. This game-theoretic perspective captures interdependent behavior by modeling other agents as minimizing their own cost as a function of the decisions of all players in the scene.

III. FORMALIZING CONTINGENCY GAMES

In this letter, we consider settings where a game-theoretic interaction model is not fully specified. For example, an autonomous car may not be sure if a nearby pedestrian intends to jaywalk; an assistive robot may not know which tool a surgeon will wish to use next. In such instances, the robot (agent R) can construct a dynamic game in which components of the model depend upon an unknown parameter $\theta \in \Theta$, with $|\Theta| = K < \infty$. When the robot has some prior information—e.g., from observations of past behavior—it can maintain a probability distribution or *belief*, $b(\theta)$, over the set of possible games. Naturally, however, this belief changes as a function of

the human's behavior *and* the robots actions. It is this *dynamic* nature of the uncertainty that the robot can exploit to come up with more efficient plans. In the formal description below, we adopt a two-player convention for clarity. We note, however, that the formalism can incorporate more players as illustrated in Section VII-B.

Approximations and modeling assumptions: Contingency games approximate a game which exactly models dynamic uncertainty, which are generally intractable to solve. Specifically, we introduce the following key modeling assumptions to facilitate tractable online computation:

- 1) We assume unilateral uncertainty with discrete support, i.e. $|\Theta| = K$. That is, while the robot R has uncertainty in the form of a discrete probability mass function $b(\theta)$, the human H acts rationally under the true hypothesis $\hat{\theta}$.
- 2) We assume the robot has access to (an estimate of) the so-called branching time, t_b , which models the future time at which additional state observations will have resolved all uncertainty about the initially unknown θ .
- 3) We simplify the robot's belief dynamics. Instead of capturing the exact belief dynamics across the planning horizon, the robot distinguishes two phases: before t_b the belief is fixed at $b(\cdot)$; after t_b the belief collapses to certainty about a single hypothesis.

We envision contingency games to be employed in a receding horizon (MPGP) fashion. In that context, beliefs are updated between each planner invocation and the branching time may vary and can be estimated online.¹ We discuss considerations and results for this case in Sections VI and VIII.

Notation conventions: We consider interaction of agents $i \in \{R, H\}$ over $T < \infty$ time steps. In contrast to existing game-theoretic formulations [12], [13], [29], [30], in a contingency game we endow each player with multiple trajectories; one for each hypothesis θ . At each $t \in [T] = \{1, 2, \dots, T\}$ and hypothesis $\theta \in \Theta$ the *state* of the game is comprised of separate variables for each agent i , i.e., $x_{\theta,t} := (x_{\theta,t}^R, x_{\theta,t}^H)$. Each agent i begins at a fixed initial state $\hat{x} := (\hat{x}^R, \hat{x}^H)$, i.e., $x_{\theta,1}^i = \hat{x}^i$, which evolves over time as a function of that agent's control action, $u_{\theta,t}^i$. States and control actions are assumed to be real vectors of arbitrary, finite dimension. For brevity, we introduce the following shorthand: we use $z_{\theta,t}^i := (x_{\theta,t}^i, u_{\theta,t}^i)$ to denote a state-control tuple, we use boldface to denote aggregation over time, e.g. $\mathbf{z}_{\theta}^i := (z_{\theta,t}^i)_{t \in [T]}$, we omit player indices to denote aggregation, e.g. $\mathbf{z}_{\theta} = (\mathbf{z}_{\theta}^R, \mathbf{z}_{\theta}^H)$, and we denote the finite collection of all $K = |\Theta|$ trajectories for player i as $\mathbf{z}_{\Theta}^i = (z_{\Theta,t}^i)_{t \in [T]}$.

Contingency game formulation: With these conventions in place, we formulate a contingency game as follows. The robot wishes to optimize its expected performance over all hypotheses (1a) while restricting all contingency plans to be feasible with respect to hypothesis-dependent constraints h_{θ} (1b) and enforcing the *contingency constraint* (1c) that the first $t_b - 1$ control inputs must be identical across all hypotheses $\theta \in \Theta$:

$$R : \mathcal{S}^R(\mathbf{z}_{\Theta}^H) := \arg \min_{\mathbf{z}_{\Theta}^R} \sum_{\theta \in \Theta} b(\theta) J^R(\mathbf{z}_{\theta}^R, \mathbf{z}_{\theta}^H) \quad (1a)$$

$$h_{\theta}^R(\mathbf{z}_{\theta}^R, \mathbf{z}_{\theta}^H) \geq 0, \quad \forall \theta \in \Theta \quad (1b)$$

¹Note that, despite the use of assumption 3 within our game formulation, we will employ a belief updater that still captures more accurate belief dynamics as we shall discuss in Section VI.

$$c(\mathbf{u}_{\Theta}^R; t_b) = 0. \quad (1c)$$

Simultaneously, the robot interacts with K versions of agent H, each of which is guided by a different intent θ which parameterizes both the hypothesis-dependent cost J_{θ}^H (2a) and constraints h_{θ}^H (2b):

$$\forall \theta \in \Theta : \begin{cases} H_{\theta} : \mathcal{S}^{H_{\theta}}(\mathbf{z}_{\theta}^R) \\ := \arg \min_{\mathbf{z}_{\theta}^H} J_{\theta}^H(\mathbf{z}_{\theta}^R, \mathbf{z}_{\theta}^H) \end{cases} \quad (2a)$$

$$h_{\theta}^H(\mathbf{z}_{\theta}^R, \mathbf{z}_{\theta}^H) \geq 0. \quad (2b)$$

In this contingency game formulation, it is important to appreciate that the robot's strategy depends on the distribution $b(\theta)$ and *all* the trajectories $\mathbf{z}_{\Theta} = (\mathbf{z}_{\theta})_{\theta \in \Theta}$, while H_{θ} 's strategy depends *only* on the trajectories under hypothesis θ , $\mathbf{z}_{\theta}$. Since the objectives of R and H may in general conflict with one another, and we model agents as being self-interested, such a game is termed *noncooperative*.

Taken together, optimization problems (1), (2) take the form of a generalized Nash equilibrium problem (GNEP). Solutions to this problem are defined as follows.

Definition 1: (Generalized Nash Equilibrium, [31, Ch. 1]). A trajectory profile $(\mathbf{z}_{\Theta}^{R*}, \mathbf{z}_{\Theta}^{H*})$ is a *generalized Nash equilibrium* (GNE) of the game from (1), (2) if and only if

$$\mathbf{z}_{\Theta}^{R*} \in \mathcal{S}^R(\mathbf{z}_{\Theta}^{H*}) \quad \text{and} \quad \bigwedge_{\theta \in \Theta} \mathbf{z}_{\theta}^{H*} \in \mathcal{S}^{H_{\theta}}(\mathbf{z}_{\theta}^{R*}).$$

In practice, we relax Definition 1 and only require *local* minimizers of (1), (2). Under appropriate technical qualifications, such local equilibria are characterized by first and second order conditions commensurate with concepts in optimization. Readers are directed to [31] for further details.

A solution method for contingency games is discussed in Section V; but first, we take a step back to build intuition about the behavior induced by the proposed interaction model.

IV. FEATURES OF CONTINGENCY GAME SOLUTIONS

Scenario 1. driving near a jaywalking pedestrian: Consider the scenario shown in Fig. 2. Here, the robot (R) wants to move forward but does not know if the pedestrian (H) wants to reach the left or right side of the road. To model this ambiguity, let $\theta \in \Theta := \{\text{left}, \text{right}\}$. Robot R plans a trajectory for each hypothesis, i.e. $\mathbf{z}_{\Theta}^R := (\mathbf{z}_{\text{left}}^R, \mathbf{z}_{\text{right}}^R)$, each of which is tailored to react appropriately to human H_{θ} in the corresponding scenario. Similarly, H maintains $\mathbf{z}_{\Theta}^H := (\mathbf{z}_{\text{left}}^H, \mathbf{z}_{\text{right}}^H)$, which capture its "ground truth" behavior under each hypothesis. Each of H's trajectories will react to the corresponding trajectory of R. We model the robot as a kinematic unicycle and the pedestrian as a planar point mass. To ensure collision avoidance, the robot must pass behind the pedestrian, i.e. not between the pedestrian and its (initially unknown) goal position. We capture all of these constraints via a single vector-valued function $h_{\theta}^i(\mathbf{z}_{\theta}^R, \mathbf{z}_{\theta}^H) \geq 0$, which explicitly depends on hypothesis θ .

How each agent acts in a contingency game formulation of this problem is largely affected by two quantities: (i) the initial belief that the robot holds over the human's intent and (ii) the branching time which models the time at which the robot will get certainty. We discuss the role of both quantities below.

Qualitative behavior. Role of belief: First, we keep the branching time fixed and analyze the contingency plans for a suite of beliefs. Fig. 2(a)–(e) show how both the robot's contingency

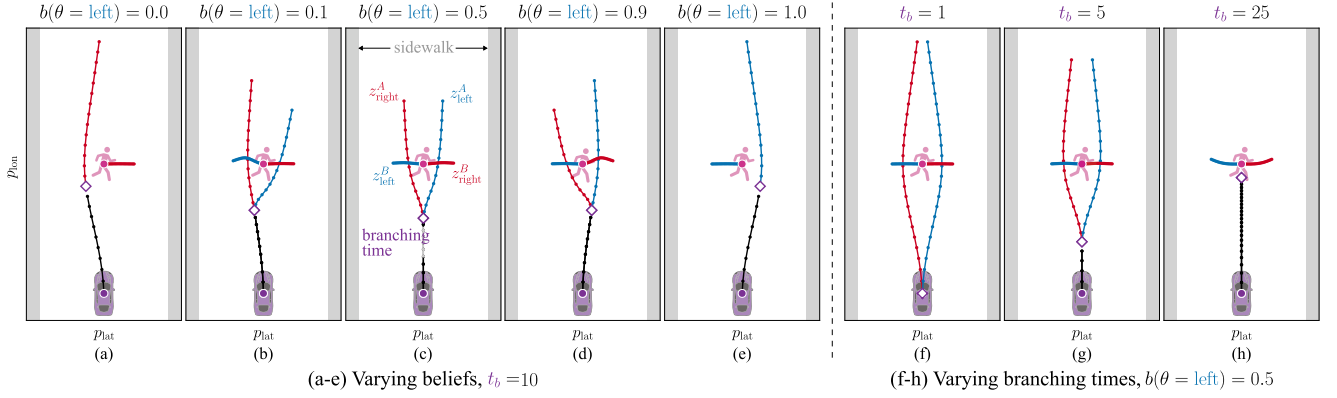


Fig. 2. Contingency game solutions for the jaywalking scenario at varying beliefs and branching times.

plan and the pedestrian's reaction change as a function of the robot's intent uncertainty. At extreme values of the belief, the robot is certain which hypothesis is accurate and the contingency strategy is equivalent to that of a single-hypothesis game with intent certainty. At intermediate beliefs, the contingency game balances hypotheses' cost according to their likelihood, yielding interesting behaviors: in Fig. 2(d) the robot plans to inch forward at first (black), but biases its initial motion towards the scenario where the pedestrian will go left (blue). Nevertheless, it still generates a plan for the less likely event that the pedestrian goes right (red).

Qualitative behavior. Role of branching time: Next, we assume that the robot's belief is always a uniform distribution, and vary the contingency game's t_b parameter. Fig. 2(f)–(h) show how extreme values of this parameter automatically recover existing variants of game-theoretic planning under uncertainty: certainty-equivalent games [1], [2], [3] and non-contingent games that plan in expectation [4], [6]. At one extreme, when $t_b = 1$, the contingency constraint (1c) is removed entirely and we obtain the solutions of the fully observed game under each hypothesis (see Fig. 2(f)). These are precisely the solutions found for the certainty-equivalent games in Fig. 2(a), (e). Note, however, that in this special case the contingency plan is not immediately actionable since it first requires the ego agent to commit to a single branch, e.g., by considering only the most likely hypothesis [1], [2], [3]. While easy to implement, such an approach can be overly optimistic and can lead to unsafe behavior since the control sequence from the selected branch may be infeasible for another intent hypothesis. For example, if the robot were to commit to $\theta = \text{left}$ in Fig. 2(f), it would be on a collision course with 50% probability.

By contrast, at the upper extreme of the branching time, $t_b = T$, the contingency plan no longer branches and it consists of a single control sequence for the entire planning horizon, cf. [4], [6]. As a result, this plan is substantially more conservative since it must trade off the cost under *all* hypotheses while respecting the constraints for any hypothesis with non-zero probability mass: in Fig. 2(g), the ego agent plans to slow down aggressively since its uncertainty about the pedestrian's intent renders both sides of the roadway blocked.

Overall, this analysis shows that contingency games (i) unify various popular approaches for game-theoretic planning under uncertainty, and (ii) go beyond these existing formulations towards games that consider the *distribution* of other players' intents while *anticipating intent certainty* in the future.

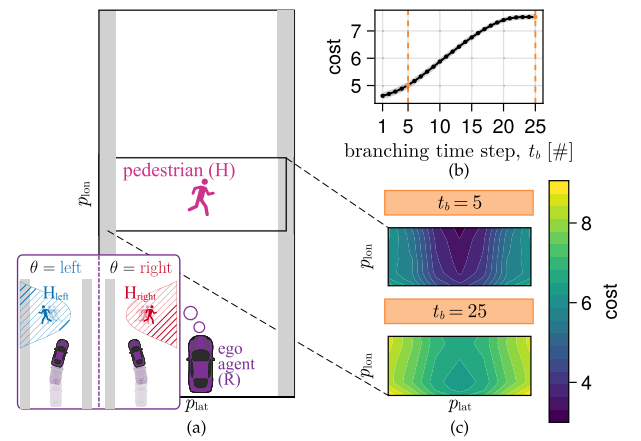


Fig. 3. Cost of contingency plans in the jaywalking scenario. (a) State sampling region. (b) Cost per t_b averaged over all initial states. (c) Spatial cost distribution for two fixed t_b .

Quantitative impact of the branching time: Given the central role of the branching time in our interaction model, we further analyze the quantitative impact of this parameter on the robot's contingency plan. For this purpose, we generate contingency plans across varying branching times, $t_b \in \{1, \dots, 25\}$, for each of 70 different initial pedestrian positions sampled from a uniform grid around the nominal position as shown in Fig. 3(a). Fig. 3(b) shows that, by anticipating future certainty at earlier branching times ($t_b = 5$), the robot discovers lower-cost plans than a method that assumes uncertainty will never resolve ($t_b = 25$). Furthermore, the spatial distribution of the cost in Fig. 3(c) reveals that robot generates particularly low-cost plans for $t_b = 5$ if the human is initially in the center of the road. Here, the contingency plan exploits the fact that—irrespective of the human's true intent—the road will be cleared by the time the robot arrives, cf. Fig. 2(g). Of course, this analysis pertains to the *open-loop* plan. However, as we shall demonstrate in Section VIII, a performance advantage persists under the added effect of receding-horizon planning.

V. TRANSFORMING CONTINGENCY GAMES INTO MIXED COMPLEMENTARITY PROBLEMS

Next, we discuss how to compute strategies from this interaction model. Rather than developing a specialized solver, we

demonstrate how contingency games can be transformed into mixed complementarity problems (MCPs) for which large-scale off-the-shelf solvers are readily available [7]. Our implementation of a game-theoretic contingency planner internally synthesizes such MCPs from user-provided descriptions of dynamics, costs, constraints, and beliefs.

We begin by deriving the Karush-Kuhn-Tucker (KKT) conditions for the contingency game in (1), (2). Under an appropriate constraint qualification (e.g., Abadie, or linear independence [31, Ch. 1], [32, Ch. 12]), these first-order conditions are jointly necessary for any generalized Nash equilibrium of the game. The Lagrangians of both players are:

$$\begin{aligned} R : \mathcal{L}^R(z_\Theta^R, z_\Theta^H, \lambda_\Theta^R, \rho) &= \rho^\top c(u_\Theta^R; t_b) \\ &+ \sum_{\theta \in \Theta} (b(\theta) J^R(z_\theta^R, z_\theta^H) - \lambda_\theta^{\top} h_\theta^R(z_\theta^R, z_\theta^H)), \\ H_\theta : \mathcal{L}_\theta^H(z_\theta^R, z_\theta^H, \lambda_\theta^H) &= J_\theta^H(z_\theta^R, z_\theta^H) - \lambda_\theta^{\top} h_\theta^H(z_\theta^R, z_\theta^H), \end{aligned} \quad (3)$$

where ρ is the Lagrange multiplier for player R's contingency constraint, and λ_θ^i are Lagrange multipliers for all other constraints of player i at hypothesis θ . Denoting complementarity by $\perp$, we derive the following KKT system for all players:

$$\forall \theta \in \Theta : \begin{cases} \nabla_{z_\theta^R} \mathcal{L}^R(z_\theta^R, z_\theta^H, \lambda_\theta^R, \rho) = 0, & (4a) \\ \nabla_{z_\theta^H} \mathcal{L}_\theta^H(z_\theta^R, z_\theta^H, \lambda_\theta^H) = 0, & (4b) \\ 0 \leq h_\theta^R(z_\theta^R, z_\theta^H) \perp \lambda_\theta^R \geq 0, & (4c) \\ 0 \leq h_\theta^H(z_\theta^R, z_\theta^H) \perp \lambda_\theta^H \geq 0, & (4d) \\ c(u_\Theta^R; t_b) = 0. & (4e) \end{cases}$$

Collectively, (4) forms an MCP, as defined below [31, Ch. 1].

Definition 2 (Mixed Complementarity Problem): A mixed complementarity problem (MCP) takes the following form: Given $G : \mathbb{R}^d \rightarrow \mathbb{R}^d$, lower bounds $v_{lo} \in [-\infty, \infty]^d$ and upper bounds $v_{up} \in (-\infty, \infty]^d$, solve for $v^* \in \mathbb{R}^d$ such that, for each $j \in \{1, \dots, d\}$, one of the equations below holds:

$$\begin{aligned} [v]_j^* &= [v]_{lo,j}, [G]_j(v^*) \geq 0, \\ [v]_{lo,j} &< [v]_j^* < [v]_{up,j}, [G]_j(v^*) = 0, \\ [v]_j^* &= [v]_{up,j}, [G]_j(v^*) \leq 0, \end{aligned}$$

where $[G]_j$ denotes the j^{th} component of G , and $[v]_{lo,j}$ and $[v]_{up,j}$ denote the j^{th} components of v_{lo} and v_{up} , respectively.

Observe that the KKT conditions (4) encode an MCP with variable v , function G , and bounds v_{lo} , v_{up} block-wise defined (by slight abuse of notation): one block for each $\theta \in \Theta$,

$$[v]_\theta = (z_\theta^R, z_\theta^H, \lambda_\theta^R, \lambda_\theta^H), \quad (5a)$$

$$[v]_{lo,\theta} = (-\infty, -\infty, 0, 0), \quad (5b)$$

$$[v]_{up,\theta} = (\infty, \infty, \infty, \infty), \quad (5c)$$

$$[G(v)]_\theta = \begin{bmatrix} \nabla_{z_\theta^R} \mathcal{L}^R(z_\theta^R, z_\theta^H, \lambda_\theta^R, \rho) \\ \nabla_{z_\theta^H} \mathcal{L}_\theta^H(z_\theta^R, z_\theta^H, \lambda_\theta^H) \\ h_\theta^R(z_\theta^R, z_\theta^H) \\ h_\theta^H(z_\theta^R, z_\theta^H) \end{bmatrix}, \quad (5d)$$

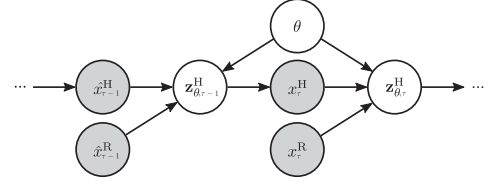


Fig. 4. Bayesian network modeling the robot's intent inference problem. Shaded nodes represent observed variables.

and an additional block for the contingency constraint

$$[v]_c = \rho, [v]_{lo,c} = -\infty, [v]_{up,c} = \infty, [G(v)]_c = c(u_\Theta^R; t_b). \quad (6)$$

To establish that the MCP solution is indeed a local equilibrium of the contingency game, sufficient second-order conditions [32, Thm. 12.6] can be checked for (1) and (2).

VI. ONLINE PLANNING WITH CONTINGENCY GAMES

We envision contingency games to be deployed in a MPGP framework where beliefs and branching times are estimated online. Next, we discuss considerations for this setting.

A. Belief Updates

While the true human intent $\hat{\theta}$ is hidden from the robot, the robot can utilize observations of past human decisions to update its current belief $b_\tau(\theta) = P(\theta | \hat{x})$ about this quantity, where $\hat{x} := \{(\hat{x}_1^R, \hat{x}_1^H), \dots, (\hat{x}_\tau^R, \hat{x}_\tau^H)\}$ is the sequence of joint human-robot states observed up until the current time τ . We use the model in Fig. 4 to cast this inference problem in a Bayesian framework. In this model, we use our game-theoretic planner to compute jointly a *nominal* state-action trajectory for each agent and for each hypothesis. The robot's portion of this solution computed at time τ , $z_{\theta,\tau}^R$, serves as our receding-horizon motion plan, and the human's portion of the plan, $z_{\theta,\tau}^H$, constitutes a nominal receding-horizon prediction of their future actions under each hypothesis $\theta \in \Theta$. However, due to bounded rationality [33], the human may not execute exactly this plan. Hence, at the next time step $\tau + 1$ when we observe a new human physical state, $\hat{x}_{\tau+1}^H$, we treat it as a random emission of the *previous* nominal predictions. Similar to prior works [22], [24], in our experiments we assume that human states are distributed according to a Gaussian mixture model with one mode for each hypothesis θ ; i.e., $p(\hat{x}_{\tau+1}^H | z_{\theta,\tau}^H) = \mathcal{N}(\mu_\theta, \Sigma)$ where the mean μ_θ is the expected human state extracted from the *previous* human prediction, $z_{\theta,\tau}^H$, and Σ is a covariance parameter characterizing human rationality. In summary, the robot can recursively update their belief via

$$b_{\tau+1}(\theta) = \frac{p(\hat{x}_{\tau+1}^H | z_{\theta,\tau}^H) b_\tau(\theta)}{\sum_{\theta' \in \Theta} p(\hat{x}_{\tau+1}^H | z_{\theta',\tau}^H) b_\tau(\theta')}. \quad (7)$$

B. Estimating Branching Times

Prior work on non-game-theoretic contingency planning either chooses the branching time as informed by a careful heuristic design [19] or through offline analysis of the belief updater [34]. A thorough theoretical analysis on how to choose the branching time t_b in the more general game-theoretic case is beyond the scope of this work. Nonetheless, to demonstrate

the utility of anticipating future intent certainty, we propose a branching time heuristic which only requires access to the previous game solution and current belief.

Heuristic branching time estimator: Intuitively, the branching time should be lower when the future human actions are more “distinct” under each hypothesis. Our heuristic captures this relationship as follows. Let $\mathcal{H}[b_\tau] = -\sum_{\theta \in \Theta} b_\tau(\theta) \log_{|\Theta|}(b_\tau(\theta))$ denote the entropy of belief b_τ . Furthermore, let $B(z_{\Theta, \tau-1}^H, \theta, k)[\cdot]$ denote the operator that, for a given θ , takes the first k states from the previously-computed $z_{\Theta, \tau-1}^H$ as hypothetical observations and returns the thus updated belief. We approximate the branching time as

$$t_b(b_\tau, z_{\Theta, \tau-1}^H) = \max_{\theta \in \Theta} \min_{k \in \{2, \dots, T\}} k \quad (8)$$

s.t. $\mathcal{H}[B(z_{\Theta, \tau-1}^H, \theta, k)[b_\tau]] \leq \epsilon$.

Note that the minimum branching time chosen by this heuristic is $t_b = 2$, ensuring that the robot has a unique first input to apply during receding-horizon operation. Procedurally, this heuristic is straightforward to implement: for *each* hypothesis, we predict the belief via a hypothetical observation sequence as if the human (i) is perfectly rational and (ii) does not re-plan in the future; then, we return the first time at which *all* predicted beliefs reach entropy threshold ϵ .² Assumptions (i) and (ii) make this heuristic cheap to evaluate since they avoid the need to re-compute game solutions within (8). While, these approximations may affect the accuracy of the estimator we shall demonstrate the utility of this approach in Section VIII.

VII. EXPERIMENTAL SETUP

We wish to study the value of *anticipating* future certainty in game-theoretic planning. Therefore, we compare our method against a non-contingent game-theoretic baselines on two simulated interaction scenarios in which dynamic uncertainty naturally occurs.

A. Compared Methods

Beyond the contingency game that uses our branching time heuristic, we consider the following methods, all of which operate in receding-horizon fashion with online belief updates according to (7). Note that, following the discussion in Section IV, all game-theoretic methods below can be understood as a contingency game with a special branching time choice.

Baseline 1. Certainty-equivalent ($t_b = 1$): This baseline assumes certainty, making a point estimate by considering only the most probable hypothesis at each time step.

Baseline 2. Fixed uncertainty ($t_b = T$): Similar to [4], this baseline ignores future information gains, assuming fixed uncertainty along the entire planning horizon.

Baseline 3. MPC: This baseline uses non-game-theoretic model-predictive control (MPC), forecasting opponent trajectories assuming constant ground-truth velocity.

Contingency game with $t_b = 2$: To test the utility of our branching time heuristic, we also consider a contingency game that assumes certainty one step into the future—an assumption also used in non-game-theoretic contingency planning [21].

²A low threshold results in more conservative behavior. As informed by a parameter sweep over ϵ , we choose $\epsilon = 2^{-2}$ for all experiments for best performance.

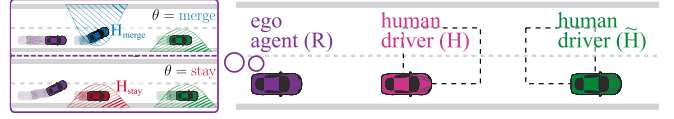


Fig. 5. Scenario 2: an autonomous vehicle seeks to overtake slow traffic on a highway while being uncertain about the lane changing intentions of the vehicle ahead.

Contingency game with oracle branching time: Additionally, we consider an oracle branching time estimator that recovers the true branching time by first simulating receding-horizon interaction with a nominal branching time and then extracting the time of certainty from the belief evolution in *hindsight*. Naturally, this oracle requires access to the *true* human intent and hence is not realizable in practice. Nonetheless, we include this variant to demonstrate the potential performance achievable with our interaction model.

B. Driving Scenarios

Beyond the jaywalking example (Scenario 1) introduced in Section IV, we evaluate our method on the following three-player scenario.

Scenario 2. highway overtaking: In this scenario, an autonomous vehicle attempts to overtake a human-operated vehicle with additional slow traffic in front, cf. Fig. 5. To perform this overtaking maneuver safely, the robot must reason about possible lane changing intentions of the other vehicle. Since the robot is uncertain about the target lane of human-driven car in front, it maintains a belief over the hypothesis space $\Theta = \{\text{merge}, \text{stay}\}$. We add a non-convex collision avoidance constraint between pairs of players which enforces that cars cannot be overtaken on the side of their target lane.

Implementation details: Throughout all experiments, we model cars as kinematic unicycles, and pedestrians as planar point masses. All systems evolve in discrete-time with time discretization of 0.2 s and agents plan over a horizon of $T = 25$ time steps. In both scenarios, road users’ costs comprise of a quadratic control penalty and their intent $\theta \in \Theta$ dictates a goal position for pedestrians and a reference lane for cars via a quadratic state cost. Collision avoidance constraints are shared between all agents.

VIII. SIMULATED INTERACTION RESULTS

The following evaluations are designed to support the claims that (C1) contingency games close the gap between the two extremes shown in Fig. 1: providing more efficient plans than fixed-uncertainty games at higher levels of safety than certainty-equivalent games; and that (C2) our branching time heuristic improves the performance of contingency games over a naive fixed branching time estimate of $t_b = 2$.

Data collection: We evaluate all methods in a large-scale Monte Carlo study as follows. For each scenario, we simulate receding-horizon interactions of 6s duration. As in the open-loop evaluation of Section IV, we repeat the simulation for 70 initial states of each non-ego agent. These initial states are drawn from a uniform grid over the state regions shown in Figs. 3(a) and 5. We generate the ground-truth human behavior from a game solution at each fixed hypothesis $\hat{\theta} \in \Theta$. Since this true human intent is initially unknown to the robot, the robot starts with a uniform

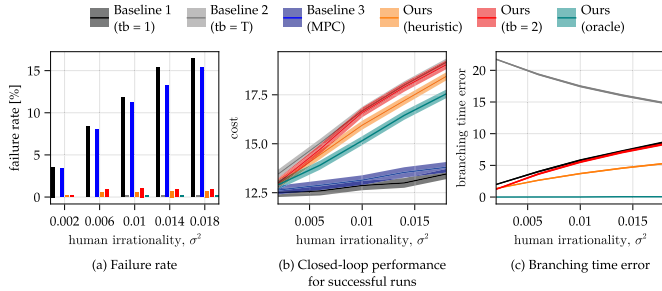


Fig. 6. Quantitative closed-loop results for the jaywalking example.

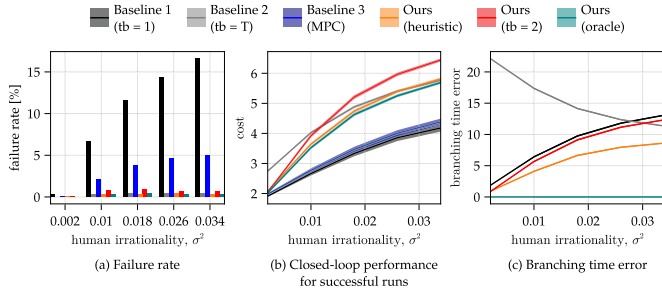


Fig. 7. Quantitative closed-loop results for the overtaking example.

belief. To test the methods under varying information gain dynamics—and thereby varying branching times—we consider five levels of human rationality, σ^2 , parameterizing an isotropic observation model, i.e., $\Sigma = \sigma^2 \mathbf{I}$ where $\mathbf{I}$ denotes the identity matrix. In contrast to the open-loop evaluation of Section IV, here the robot re-plans at every time step with the latest online estimate of the belief and branching time. Figs. 6, 7 summarize the results of this Monte Carlo study.

Quantitative results: In terms of safety, Figs. 6(a) and 7(a) show that the methods making a single prediction about the future, Baseline 1 and Baseline 3, fail significantly more often than the remaining approaches, all of which achieve failure rates below 1% across all levels of human rationality. In terms of efficiency, Figs. 6(b) and 7(b) show that contingency games incur a lower interaction cost than the more conservative fixed-uncertainty Baseline 2. However, this performance advantage relies on a dynamic branching time estimate, as indicated by the performance advantage of Ours (heuristic) over Ours ($t_b = 2$). Finally, Ours (oracle) further improves efficiency due to the tight branching time estimate (cf. Figs. 6(c) and 7(c)), demonstrating the potential of our interaction model. In summary, these results support our claims **C1** and **C2** above.

Qualitative results: To further contextualize these results with respect to claim **C1**, we visualize examples of the closed-loop behavior generated by our method, the certainty-equivalent Baseline 1, and the fixed-uncertainty Baseline 2 in Figs. 8, 9. Here, we show both the sequence of states traced out by all players and the robot's plan five time steps into the interaction. In both scenarios, the robot's initial observations cause its belief to favor an incorrect hypothesis. This belief prompts Baseline 1 to commit to an unsafe strategy, causing a collision with the human. Baseline 2, on the other hand, brakes conservatively in the face of this belief and only accelerates once the uncertainty fully resolves. Finally, our contingency game planner *anticipates*

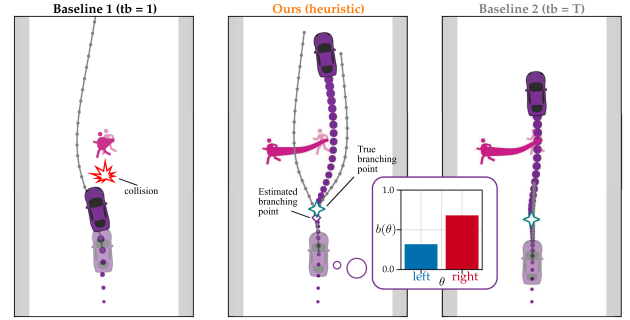


Fig. 8. Qualitative closed-loop results for the jaywalking example.

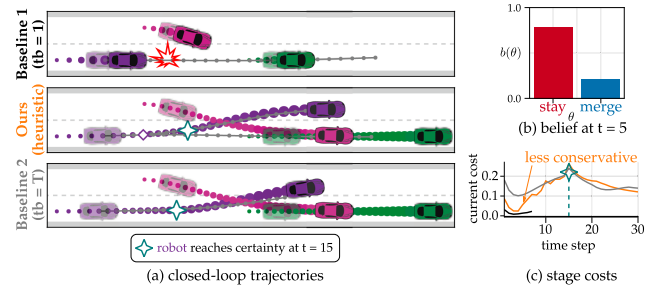


Fig. 9. Qualitative closed-loop results for the overtaking example.

the future information gain and avoids excessive braking while remaining safe.

IX. LIMITATIONS & FUTURE WORK

Even with a simple branching time heuristic, our approach outperforms the baselines. Nonetheless, the observed performance gains for the branching time oracle motivate further research into more precise estimators. Beyond that, in this work we assumed access to a fixed set of suitable intent hypotheses and our approach relies on receding-horizon re-planning to adapt to changes of this set. Future work should seek to automate the discovery of intent hypotheses, test our approach in scenarios with more complex intent dynamics, and consider extensions that explicitly capture these effects in the contingency game. Furthermore, the complexity of our approach is proportional to the product of the number of individual intents of each player, and the technique we employ to solve these games generally scales cubically with regards to total strategy size. Future work could consider employing a learning-based predictor [35] to automatically identify high-likelihood intents in complex scenarios and sub-select local players [36]. Finally, future work may extend our approach to continuous hypothesis spaces, e.g., by sampling as in [5], [37].

X. CONCLUSION

We present contingency games, a game-theoretic motion planning framework that enables a robot to efficiently interact with other agents in the face of uncertainty about their intents. By capturing simplified belief dynamics, our method allows a robot to anticipate future changes in its belief during strategic interactions. In detailed simulated driving experiments, we characterized both the qualitative behavior induced by contingency

games and the quantitative performance gains with respect to non-contingent baselines.

ACKNOWLEDGMENT

Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Research Council Executive Agency. Neither the European Union nor the granting authority can be held responsible for them.

REFERENCES

- [1] D. Sadigh, S. Sastry, S. A. Seshia, and A. D. Dragan, "Planning for autonomous cars that leverage effects on human actions," in *Proc. Robot. Sci. Syst.*, 2016, pp. 1–9.
- [2] N. Mehr, M. Wang, M. Bhatt, and M. Schwager, "Maximum-entropy multi-agent dynamic games: Forward and inverse solutions," *IEEE Trans. Robot.*, vol. 39, no. 3, pp. 1801–1815, Jun. 2023.
- [3] X. Liu, L. Peters, and J. Alonso-Mora, "Learning to play trajectory games against opponents with unknown objectives," *IEEE Robot. Automat. Lett.*, vol. 8, no. 7, pp. 4139–4146, Jul. 2023.
- [4] F. Laine, D. Fridovich-Keil, C.-Y. Chiu, and C. Tomlin, "Multi-hypothesis interactions in game-theoretic motion planning," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2021, pp. 8016–8023.
- [5] S. L. Cleac'h, M. Schwager, and Z. Manchester, "LUCIDGames: Online unscented inverse dynamic games for adaptive trajectory prediction and planning," *IEEE Robot. Automat. Lett.*, vol. 6, no. 3, pp. 5485–5492, Jul. 2021.
- [6] R. Tian, L. Sun, A. Bajcsy, M. Tomizuka, and A. D. Dragan, "Safety assurances for human-robot interaction via confidence-aware game-theoretic human models," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2022, pp. 11229–11235.
- [7] S. P. Dirkse and M. C. Ferris, "The path solver: A non-monotone stabilization scheme for mixed complementarity problems," *Optim. Methods Softw.*, vol. 2, pp. 123–156, 1995.
- [8] J. F. Fisac, E. Bronstein, E. Stefansson, D. Sadigh, S. S. Sastry, and A. D. Dragan, "Hierarchical game-theoretic planning for autonomous vehicles," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2019, pp. 9590–9596.
- [9] O. So, P. Drews, T. Balch, V. Dimitrov, G. Rosman, and E. A. Theodorou, "MPOGames: Efficient multimodal partially observable dynamic games," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2023, pp. 3189–3196.
- [10] R. Spica, E. Cristofalo, Z. Wang, E. Montijano, and M. Schwager, "A real-time game theoretic planner for autonomous two-player drone racing," *IEEE Trans. Robot.*, vol. 36, no. 5, pp. 1389–1403, Oct. 2020.
- [11] S. Musić and S. Hirche, "Haptic shared control for human-robot collaboration: A game-theoretical approach," *IFAC-LettersOnLine*, vol. 53, pp. 10216–10222, 2020.
- [12] D. Fridovich-Keil, E. Ratner, L. Peters, A. D. Dragan, and C. J. Tomlin, "Efficient iterative linear-quadratic approximations for nonlinear multi-player general-sum differential games," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2020, pp. 1475–1481.
- [13] S. L. Cleac'h, M. Schwager, and Z. Manchester, "ALGAMES: A fast augmented lagrangian solver for constrained dynamic games," *Auton. Robots*, vol. 46, pp. 201–215, 2022.
- [14] E. L. Zhu and F. Borrelli, "A sequential quadratic programming approach to the solution of open-loop generalized Nash equilibria," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2023, pp. 3211–3217.
- [15] D. S. Bernstein, R. Givan, N. Immerman, and S. Zilberstein, "The complexity of decentralized control of Markov decision processes," *Math. Operations Res.*, vol. 27, pp. 819–840, 2002.
- [16] J. Goldsmith and M. Mundhenk, "Competition adds complexity," in *Proc. Conf. Neural Inf. Process. Syst.*, 2007.
- [17] W. Schwarting, A. Pierson, S. Karaman, and D. Rus, "Stochastic dynamic games in belief space," *IEEE Trans. Robot.*, vol. 37, no. 6, pp. 2157–2172, Dec. 2021.
- [18] W. Wang et al., "Social interactions for autonomous driving: A review and perspectives," *Foundations Trends Robot.*, vol. 10, no. 3–4, pp. 198–376, 2022.
- [19] J. Hardy and M. Campbell, "Contingency planning over probabilistic obstacle predictions for autonomous road vehicles," *IEEE Trans. Robot.*, vol. 29, no. 4, pp. 913–929, Aug. 2013.
- [20] W. Zhan, C. Liu, C.-Y. Chan, and M. Tomizuka, "A non-conservatively defensive strategy for urban autonomous driving," in *Proc. IEEE 19th Int. Conf. Intell. Transp. Syst.*, 2016, pp. 459–464.
- [21] Y. Chen, U. Rosolia, W. Ubellacker, N. Csomay-Shanklin, and A. D. Ames, "Interactive multi-modal motion planning with branch model predictive control," *IEEE Robot. Automat. Lett.*, vol. 7, no. 2, pp. 5365–5372, Apr. 2022.
- [22] S. H. Nair, V. Govindarajan, T. Lin, C. Meissen, H. E. Tseng, and F. Borrelli, "Stochastic mpc with multi-modal predictions for traffic intersections," in *Proc. IEEE 25th Int. Conf. Intell. Transp. Syst.*, 2022, pp. 635–640.
- [23] A. Cui, S. Casas, A. Sadat, R. Liao, and R. Urtasun, "LookOut: Diverse multi-future prediction and planning for self-driving," in *Proc. IEEE/CVF Int. Conf. Comput. Vis.*, 2021, pp. 16087–16096.
- [24] E. Tolstaya, R. Mahjourian, C. Downey, B. Vadarajan, B. Sapp, and D. Anguelov, "Identifying driver interactions via conditional behavior prediction," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2021, pp. 3473–3479.
- [25] Y. Chen, P. Karkus, B. Ivanovic, X. Weng, and M. Pavone, "Tree-structured policy planning with learned behavior models," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2023, pp. 7902–7908.
- [26] C. Packer et al., "Is anyone there? learning a planner contingent on perceptual uncertainty," in *Proc. Conf. Robot Learn.*, 2022, pp. 1607–1617.
- [27] N. Rhinehart, R. McAllister, K. Kitani, and S. Levine, "PRECOG: Prediction conditioned on goals in visual multi-agent settings," in *Proc. IEEE/CVF Int. Conf. Comput. Vis.*, 2019, pp. 2821–2830.
- [28] N. Rhinehart et al., "Contingencies from observations: Tractable contingency planning with learned behavior models," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2021, pp. 13663–13669.
- [29] T. Başar and G. J. Olsder, *Dynamic Noncooperative Game Theory*, 2nd ed. Philadelphia, PA, USA: Soc. Ind. Appl. Math., 1999.
- [30] F. Laine, D. Fridovich-Keil, C.-Y. Chiu, and C. Tomlin, "The computation of approximate generalized feedback nash equilibria," *SIAM J. Optim.*, vol. 33, no. 1, pp. 294–318, 2023.
- [31] F. Facchinei and J.-S. Pang, *Finite-Dimensional Variational Inequalities and Complementarity Problems*, vol. 1. Berlin, Germany: Springer, 2003.
- [32] J. Nocedal and S. Wright, *Numerical Optimization*. Berlin, Germany: Springer, 2006.
- [33] A. Rubinstein, *Modeling Bounded Rationality*. Cambridge, MA, USA: MIT Press, 1998.
- [34] A. Bajcsy, A. Siththaranjan, C. J. Tomlin, and A. D. Dragan, "Analyzing human models that adapt online," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2021, pp. 2754–2760.
- [35] J. Roh, C. Mavrogiannis, R. Madan, D. Fox, and S. Srinivasa, "Multimodal trajectory prediction via topological invariance for navigation at uncontrolled intersections," in *Proc. Conf. Robot Learn.*, 2021, pp. 2216–2227.
- [36] M. Chahine, R. Firoozi, W. Xiao, M. Schwager, and D. Rus, "Local non-cooperative games with principled player selection for scalable motion planning," in *Proc. IEEE/RSJ Int. Conf. Intell. Robots Syst.*, 2023, pp. 880–887.
- [37] D. Paccagnan and M. C. Campi, "The scenario approach meets uncertain game theory and variational inequalities," in *Proc. Conf. Decis. Mak. Control*, 2019, pp. 6124–6129.