

Stabilizing Welfare-Maximizing Decisions via Endogenous Transfers

Joshua Kavner

Independent Researcher

joshuakavner.research@gmail.com

Abstract

Many multi-agent systems depend on collective decisions made by self-interested agents, which raises deep questions about coalition formation and stability. We study social choice with endogenous, outcome-contingent transfers, where agents voluntarily form contracts that redistribute utility depending on the collective decision, allowing for fully strategic coalition formation. We show that under consensus rules, individually rational strong Nash equilibria (IR-SNE) always exist, implementing welfare-maximizing outcomes with feasible transfers, and provide a simple, efficient algorithm to construct them. For more general anonymous, monotonic, and resolute rules, we identify necessary conditions for viable deviations, significantly limiting the possibility of destabilizing coalitions. By bridging cooperative and noncooperative perspectives, our approach shows that transferable utility can achieve core-like stability, restoring efficiency and budget balance even where classical impossibility results generally apply. Overall, this framework offers a practical and robust way to coordinate large-scale strategic multi-agent systems.

1 Introduction

It is 2050, and the utopia of the Internet of Things has become a reality in transportation. Autonomous vehicles taxi commuters to work, transcontinental shipping is coordinated through autonomous platoons, and algorithms balance intrinsic tradeoffs between safety and efficiency. From a policy maker’s perspective, centralized solutions are appealing: by collecting global information, algorithms can guide traffic flows toward socially optimal outcomes, reduce congestion, and avoid the inefficiencies that arise when agents act myopically (e.g., Braess’s paradox [Christodoulou and Koutsoupias, 2005]). However, realizing this vision in practice is far from trivial. Computing socially optimal outcomes is often NP-hard, incentive-compatible preference elicitation is possible only under restrictive conditions, and centralized mechanisms are fragile to communication failures. As a result, purely centralized solutions struggle to deliver robust, reliable performance.

These limitations have motivated research on decentralized multi-agent systems, where self-interested agents act autonomously but coordinate through learning or negotiation [Kraus, 1997; Ephrati and Rosenschein, 1997]. For example, centralized training with decentralized execution (CTDE) allows agents to learn coordinated policies under shared information during training [Lowe *et al.*, 2017; Amato, 2024]. However, once deployed, agents cannot adapt or negotiate on the fly. In many strategic environments, agents must alternate between competition and cooperation depending on context [Axelrod and Hamilton, 1981]. Coalition formation thus provides a natural middle ground: by forming voluntary coalitions, agents can improve outcomes relative to fully decentralized behavior without requiring full centralization. Still, coalition formation comes with challenges: computing optimal coalitions is NP-hard [Sandholm *et al.*, 1999], iterative coalition deviations may fail to converge [Meir, 2017], and strategic behavior under common social choice rules remains computationally intractable [Xia *et al.*, 2009].

In this paper, we propose *endogenous transfers* as a mechanism for stabilizing collective decision-making in multi-agent systems. Collective choices among self-interested agents can naturally be modeled as a social choice problem: agents vote over alternatives, and the outcome depends on the aggregation of their preferences. Rather than prescribing coalitions or enforcing outcomes, our approach allows agents to voluntarily redistribute utility through contracts that are contingent on the chosen alternative. These transfers reshape incentives and expand the space of stable outcomes, allowing coalitions to form and dissolve strategically [Jackson and Wilkie, 2005].

Our model is inspired by classical work on bribery and coalitional manipulation in voting, as well as cooperative game-theoretic analyses of social choice [Bachrach *et al.*, 2011]. In traditional bribery or manipulation models, coalitions are often assumed to be exogenous and bribes are budget-constrained [Elkind *et al.*, 2009; Faliszewski *et al.*, 2009]. By contrast, in our framework, both coalitions and transfers are fully endogenous: agents decide how to vote and how to contract, anticipating the behavior of others while respecting individual rationality. This implies that agents can offer transfers contingent on a particular alternative winning, forming stable coalitions without any central enforcement.

The central question, then, is whether such voluntary transfers are sufficient to stabilize welfare-maximizing outcomes

in the presence of strategic, coalitional behavior. We show that under consensus rules, stable outcomes always exist and can be constructed efficiently, and we characterize fundamental limits on stability under more general voting rules.

1.1 Our Contributions

We make the following contributions.

A model of endogenous transfers in social choice. We introduce a noncooperative model in which agents may form transfer agreements in a strategic pre-round, followed by strategic voting under a fixed social choice rule. Contracts are conditioned on alternatives winning and allow agents to share the utility generated by collective decisions. This framework separates contract formation from preference revelation, enabling the study of coalition stability without assuming centralized enforcement or exogenously given coalitions.

Strong Nash equilibrium under consensus rules. For consensus rules, which generalize cooperative bargaining scenarios in which agents can veto unfavorable outcomes, we show that individually rational strong Nash equilibria (IR-SNE) always exist. Moreover, every IR-SNE induces a welfare-maximizing outcome supported by a feasible transfer scheme. We provide a simple, computationally efficient algorithm for constructing such equilibria, establishing both existence and implementability.

Necessary conditions for deviations in general. We extend our analysis beyond consensus to arbitrary anonymous, monotonic, and resolute social choice rules. In this setting, we derive a tight necessary condition for the existence of IR coalitional deviations. While this does not resolve known NP-hardness of manipulation for fixed rules, it substantially restricts the space of profitable deviations and clarifies when stability cannot be restored via transfers.

A bridge between cooperative and noncooperative social choice. Our results admit a natural interpretation in terms of cooperative game theory. Although the underlying environment is noncooperative, the presence of transferable utility allows coalitions to redistribute surplus via contracts. As a result, strong Nash equilibrium coincides with a notion of core stability with respect to the set of outcomes implementable by the social choice rule. In particular, under consensus, IR-SNE implements core-stable, welfare-maximizing allocations despite the absence of centralized mechanism design.

Finally, our framework offers a complementary perspective on classical impossibility results in mechanism design. Green and Laffont [1977] show that no mechanism can simultaneously achieve efficiency, budget balance, and strategyproofness in general quasi-linear domains [Nath and Sandholm, 2019]. Rather than contradicting this result, we refine the notion of strategic robustness. By allowing transferable utility and coalition-aware contracts, we replace unconditional truthfulness with a coalition-proof requirement: no group of agents can profitably deviate given feasible transfers among its members. In this sense, IR-SNE can be viewed as a refinement of strategyproofness that restores efficiency and budget balance within the space of coalitionally stable outcomes.

1.2 Related Work

The computational complexity of bribery in elections was first studied by Faliszewski *et al.* [2006], who asked whether a preferred candidate can be made the winner under a voting rule by modifying the preferences of a limited number of agents. Subsequent work analyzed bribery under many voting rules, including Borda [Brelsford *et al.*, 2008], maximin [Faliszewski *et al.*, 2011], Schulze [Parkes and Xia, 2012], k -approval [Lin, 2012], STV and ranked pairs [Xia, 2012], and Bucklin and fallback voting [Faliszewski *et al.*, 2015]. Faliszewski [2008] introduced non-uniform bribery prices per agent under a fixed budget, while Faliszewski *et al.* [2007] and Elkind *et al.* [2009] studied models where the cost of a bribe depends on the swap distance between the original and modified rankings. For comprehensive overviews, see Baumeister and Rothe [2016] and Faliszewski and Rothe [2016].

Bribery is closely related to coalitional manipulation, which studies whether a preferred outcome can be achieved by coordinated misreports from a coalition of agents [Conitzer *et al.*, 2007; Xia *et al.*, 2009]. Unlike bribery, coalitional manipulation typically assumes the coalition is exogenously given; endogenously selecting coalitions and distributing gains is itself computationally intractable [Sandholm *et al.*, 1999; Chalkiadakis *et al.*, 2011]. Bachrach *et al.* [2009] introduced the cost of stability, and Bachrach *et al.* [2011] showed that bribery can be interpreted as utility-constrained coalition recruitment, linking voting with cooperative game theory. Our model extends this perspective by explicitly separating preference reports from transfer contracts and by imposing individual rationality on every participating agent, leading to the notion of IR-SNE. This allows us to construct feasible, budget-balanced transfer schemes that stabilize outcomes under consensus and related rules, complementing Bachrach *et al.*'s existential and complexity results with a constructive, mechanism-aware analysis.

Related questions about side payments and commitments have been studied in game theory [Fershtman *et al.*, 1991; Jackson and Wilkie, 2005; Monderer and Tennenholtz, 2009; Kalai and Kalai, 2013; Turrini, 2013; Grandi *et al.*, 2019; Geffner *et al.*, 2025] and in multi-agent RL [Sodomka *et al.*, 2013; Yang *et al.*, 2020; Haupt *et al.*, 2024; Kolumbus *et al.*, 2024]. A separate literature examines vote trading through competitive equilibrium models [Casella *et al.*, 2012; Casella and Palfrey, 2021; Casella and Macé, 2021].

2 Preliminaries

Model. An instance $I = \langle \mathcal{N}, \mathcal{A}, u \rangle$ consists of a set of n agents $\mathcal{N}$, a set of m alternatives $\mathcal{A}$, and base utility functions $u = (u_1, \dots, u_n)$, where each $u_i : \mathcal{A} \rightarrow \mathbb{R}_{\geq 0}$. We denote the *social welfare* of an alternative a as $SW(a) = \sum_{i \in \mathcal{N}} u_i(a)$. A resolute social choice function $f : \mathcal{L}(\mathcal{A})^n \rightarrow \mathcal{A}$ selects a single alternative given votes of the agents. We assume f is anonymous, monotone, and resolute (AMR); it uses a fixed tie-breaking order.

In particular, we focus on the resolute *consensus rule*, where an alternative is chosen only if it is top-ranked by all agents; otherwise, a default alternative $a^* \in \mathcal{A}$ is selected.

Consensus is both anonymous and monotonic.

Contracts. A *contract* is a collection of transfers $c : \mathcal{N} \times \mathcal{A} \rightarrow \mathbb{R}$, where $c_{i \rightarrow j}(a)$ represents a promise from agent i to agent j conditional on alternative a . Contracts are proposed before voting and are enforced conditional on the realized winner. Each contract induces a net transfer scheme τ defined by $\tau_i(a) = \sum_{j \neq i} c_{j \rightarrow i}(a) - \sum_{j \neq i} c_{i \rightarrow j}(a)$, which satisfies budget balance: $\sum_i \tau_i(a) = 0$. We assume that the set of contracts is closed under netting, meaning that only the induced transfer scheme τ affects utilities and strategic behavior, except when the identity of transfer recipients is relevant for coalition formation. The null contract is denoted by $\emptyset$.

Strategic Voting. Given base utilities u and transfers τ , agents' realized utilities are quasi-linear: if alternative a is chosen, agent i receives $U_i(a) = u_i(a) + \tau_i(a)$. These utilities induce a truthful preference ordering $R_i^*(u + \tau) \in \mathcal{L}(\mathcal{A})$ defined by $a \succ_i^* b \iff U_i(a) > U_i(b)$, with ties broken according to a fixed ordering over $\mathcal{A}$. A vote R_i is *truthful* if $R_i = R_i^*(u + \tau)$. Agents submit votes $R_i \in \mathcal{L}(\mathcal{A})$, yielding a vote profile $P = (R_1, \dots, R_n)$. For a subset of agents $S \subseteq \mathcal{N}$, we denote by $P_S^*(u + \tau)$ the truthful vote profile of agents in S under utilities $u + \tau$. We write $P^* = P_{\mathcal{N}}^*(u)$ for the fully truthful profile under the null contract.

Coalitional Deviation. Our definition of coalitional deviation builds on concepts from *unweighted coalitional manipulation* [Xia et al., 2009] and *bribery* [Faliszewski et al., 2009], both of which assume truthful behavior by non-deviating agents. These definitions require refinement when stability is evaluated over states (P, τ, S) or along sequences of deviations, as in iterative voting [Meir, 2017]. In particular, only agents in the deviating coalition are allowed to modify outgoing transfers or vote strategically. Agents outside the coalition may experience changes in their incoming transfers due to the coalition's modifications but do not actively adjust their own outgoing contracts.

Definition 1 (Coalitional Deviation). Given base utilities u , a *coalitional deviation* from a state (P, τ, S) is any state (P', τ', S') constructed as follows. A coalition of agents $S' \subseteq \mathcal{N}$ proposes modifications to outgoing contracts $(c'_{i \rightarrow j}(a))_{i \in S', j \in \mathcal{N}, a \in \mathcal{A}}$ relative to the existing contract scheme c , inducing a new transfer scheme τ' , and submits new votes $P_{S'}$. Agents outside the coalition, $i \notin S'$, are assumed to vote truthfully with respect to their updated utilities $u + \tau'$, denoted $P_{-S'}^*(u + \tau')$. The resulting vote profile is $P' = (P_{S'}, P_{-S'}^*(u + \tau'))$, yielding outcome $f(P')$.

Under Definition 1, an agent i belongs to the deviating coalition S' if and only if they either vote strategically with respect to $u + \tau'$, or actively modify their outgoing contracts. This notion is *maximal*: agents may join a coalition in favor of an alternative even if they are not individually pivotal. Non-participating agents vote truthfully with respect to the updated utilities $u + \tau'$, modeling the independence of coalitions: each coalition chooses its deviation assuming all other agents respond truthfully to their adjusted utilities. This contrasts with an alternative, *sticky* notion, where non-participating agents neither update their votes nor their contracts. In Appendix A, we show that equilibrium may fail to exist under this sticky

variant, even in trivial settings.

Definition 2 (Individually Rational Deviation). Let (P, τ, S) be a state and (P', τ', S') a coalitional deviation. The deviation (P', τ', S') is *individually rational* (IR) if

$$U'_i(f(P')) \geq U_i(f(P)) \quad \forall i \in S',$$

with strict inequality for at least one $i \in S'$. We call (P', τ', S') *IR-feasible* from (P, τ, S) .

Definition 3 (Strong Nash Equilibrium). A state (P, τ, S) is a *strong Nash equilibrium* (SNE) if no coalition S' has an IR deviation. It is an *IR-SNE* if it is additionally IR-feasible from the truthful state $(P^*, \emptyset, \emptyset)$.

Example 1. Consider alternatives $\mathcal{A} = \{1, 2\}$, agents $\mathcal{N} = \{1, \dots, 5\}$, f as the consensus rule with default $a^* = 2$, and utilities $u = (u_i(a))_{i \in \mathcal{N}, a \in \mathcal{A}}$ defined as:

$$u = \begin{pmatrix} 2 & 1 & 2 & 1 & 11 \\ 4 & 1 & 3 & 2 & 3 \end{pmatrix}.$$

The truthful vote is $R_5^* = (1 > 2)$ and $R_i^* = (2 > 1)$, $i \in \mathcal{N} \setminus \{5\}$, so that f returns 2 by default. A feasible transfer scheme is

$$\tau = \begin{pmatrix} 2 & 0 & 1 & 1 & -4 \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix} \implies u + \tau = \begin{pmatrix} 4 & 1 & 3 & 2 & 7 \\ 4 & 1 & 3 & 2 & 3 \end{pmatrix}.$$

In $u + \tau$, all agents $i \in \mathcal{N}$ vote with $R_i = (1 > 2)$, with $P = (R_i)_{i \in \mathcal{N}}$ so that $f(P) = 1$. The set of participating agents $S = \mathcal{N}$. All agents in S weakly prefer alternative 1 in $u + \tau$ to 2 in u , so the state (P, τ, S) is IR-feasible from the truthful state $(P^*, \emptyset, \emptyset)$. It can be shown that there is no (P', τ', S') so that agents in S' weakly prefer $u_i(f(P')) + \tau'_i(f(P'))$ to $u_i(f(P)) + \tau_i(f(P))$, so (P, τ, S) is an IR-SNE.

3 Main Results

We focus on the consensus rule f , where an alternative $a \in \mathcal{A}$ is chosen if and only if all agents rank a first in their reported preferences; otherwise, a default $a^* \in \mathcal{A}$ is selected. This setting is motivated both practically and technically. Practically, consensus generalizes cooperative bargaining scenarios, in which multiple agents must determine how to split a surplus across a Pareto frontier: what benefits one agent necessarily limits another. In these settings, the default a^* represents a disagreement value that agents receive if negotiations fail [Thomson, 1994]. Consensus thus provides a concise description of real-world multi-agent planning problems in which any agent can veto a plan, forcing all agents to receive a^* [Kraus, 1997].

Technically, consensus is appealing because it allows for concise necessary and sufficient conditions for SNE, and its structure provides intuition that extends to other anonymous, monotonic, and resolute (AMR) rules. To see why, consider a state (P, τ, S) in a transferable-utility social choice game. A potential SNE can fail in exactly two ways: any coalitional deviation (P', τ', S') starting from (P, τ, S) with $f(P) = b$ either (i) changes the outcome, with $f(P') \neq b$, or (ii) preserves the outcome, with $f(P') = b$ but redistributes utilities in a way that strictly benefits the coalition.

Under consensus, the first type of deviation is tightly constrained: by Lemma 2, there is no IR-feasible state in which

the default alternative a^* becomes the winner, a consequence of the *full coverage* requirement (Definition 4) that all agents rank a non-default alternative first. The second type of deviation preserves the outcome but reallocates transfers; Lemma 3 characterizes when such redistributions are compatible with IR. Together, these lemmas pin down the structure of IR-SNE under consensus and illustrate general principles for AMR rules: welfare maximization is necessary for equilibrium (Proposition 1); outcome-changing deviations require sufficiently large reallocations of utility (Proposition 4); and outcome-preserving deviations impose stricter constraints on feasible transfers than those characterized in Lemma 3.

We now turn to the technical groundwork that allows us to formalize these ideas, starting with the role of truthful states and the observability of agent participation.

3.1 Truthful Considerations and Observable Participation

Many of our proofs focus on the contract-adjusted utilities τ , rather than the underlying individual transfers $(c_{i \rightarrow j}(a))_{i,j \in \mathcal{N}, a \in \mathcal{A}}$. We assume throughout that net transfers along transitive chains cancel out, so that an agent's utility depends only on the net transfer $\tau_i(a)$ when evaluating overall stability. This allows us to reason in terms of τ rather than tracking the full contract matrix c .

We emphasize, however, that this simplification is a sleight of hand: if the underlying contracts c are ignored entirely, so that agents only care how much they donate without regard to whom, this can break equilibrium, as shown in Appendix A. Lemma 1 is a direct corollary of this simplification: it lets us deduce whether $i \in S$ from changes in τ , but this reasoning is valid only because τ captures the net effect of transitive contracts c .

Lemma 1 (Observable Participation). *Given base utilities u , consider states (P, τ, S) and (P', τ', S') . If there exists an alternative $a \in \mathcal{A}$ such that either (1) $\tau_i(a) \leq 0$ and $\tau'_i(a) < \tau_i(a)$, or (2) $\tau_i(a) \geq 0$ and $\tau'_i(a) < 0$, then $i \in S'$.*

Proof. In both cases, agent i experiences a strictly negative change in payoff under $a \in \mathcal{A}$. Since non-participating agents of S' cannot actively modify contracts, and net transfers cancel along chains by assumption, such changes can only result from i 's own strategic actions, implying $i \in S'$. $\square$

We can therefore identify coalition members based on changes in net transfers. This reasoning allows us to analyze potential deviations in general AMR rules: if the chosen alternative $f(P)$ is not a social welfare maximizer, some coalition of agents can strictly increase their total utility by jointly reallocating transfers toward a welfare-maximizing alternative. Thus, non-welfare-maximizing states cannot be SNE.

Proposition 1 (Non-Welfare-Maximizing States Are Not SNE). *Given base utilities u and an AMR rule f , consider a state (P, τ, S) . If $f(P) \notin \arg \max_{a \in \mathcal{A}} SW(a)$, then (P, τ, S) is not an SNE.*

Proof Sketch. If $f(P)$ is not a social welfare maximizer, the grand coalition can jointly reallocate transfers toward an alternative $b \in \arg \max_{a \in \mathcal{A}} SW(a)$ such that b fully covers

$f(P)$. This deviation weakly increases total utility for the chosen alternative for all agents, showing that (P, τ, S) cannot be stable. The full proof is provided in Appendix B.2. $\square$

Consensus strengthens the connection between social welfare maximization and equilibrium stability through the notion of *full coverage*. While Proposition 1 shows that non-welfare-maximizing outcomes cannot be stable under any AMR rule, consensus imposes a stricter condition: for a non-default alternative b to win, every agent must rank b first, meaning b must fully cover the truthful winner $f(P^*)$ in post-transfer utilities. This ensures that no coalition can profitably deviate, making social welfare maximization a necessary condition for IR-SNE under consensus.

Definition 4 (Fully Covered). Given utilities u and a transfer scheme τ , an alternative $b \in \mathcal{A}$ is said to be *fully covered* by $u + \tau$ if

$$U_i(b) \geq U_i(a), \quad \forall i \in \mathcal{N}, \forall a \in \mathcal{A}.$$

Proposition 2 (Truthful State Is SNE Under Consensus). *Given base utilities u and a resolute consensus rule f , consider the truthful state $(P^*, \emptyset, \emptyset)$. If $f(P^*) \in \arg \max_{a \in \mathcal{A}} SW(a)$, then $(P^*, \emptyset, \emptyset)$ is an SNE.*

Proof Sketch. By IR and the definition of consensus, any alternative in a potential deviation must fully cover the truthful winner $f(P^*)$. Since $f(P^*)$ is a social welfare maximizer, only other welfare-maximizing alternatives can meet this requirement. Any redistribution that changes participating agents' utilities would violate IR, so no IR-feasible deviation exists. Hence, $(P^*, \emptyset, \emptyset)$ is an SNE. The full proof is provided in Appendix B.2. $\square$

3.2 IR-SNE for Consensus

Building on the discussion of truthful states and full coverage in Sec. 3.1, we turn to the construction and characterization of IR-SNE under the consensus rule f . In such equilibria, the selected alternative $f(P)$ must maximize social welfare and induce full coverage in post-transfer utilities $u + \tau$, which tightly restricts IR-feasible deviations.

Within this framework, Lemma 2 establishes that no IR-feasible deviation can make the default alternative a^* a winner, while Lemma 3 characterizes when utilities associated with a welfare-maximizing alternative can be redistributed without violating IR. Together, these results pin down the constraints governing IR-SNE under consensus.

Lemma 2. *Given base utilities u and consensus rule f , consider a state (P, τ, S) that is IR-feasible from the truthful state $(P^*, \emptyset, \emptyset)$. Suppose that $f(P^*) = a^* \notin \arg \max_{a \in \mathcal{A}} SW(a)$, that $f(P) = b \in \arg \max_{a \in \mathcal{A}} SW(a)$, and that b has full coverage under $u + \tau$. Then there is no state (P', τ', S') with $f(P') = a^*$ that is IR-feasible from (P, τ, S) .*

Proof Sketch. The argument proceeds by contradiction. Suppose that starting from a state (P, τ, S) in which the welfare-maximizing alternative b is implemented with full coverage, there exists an IR-feasible deviation to a state (P', τ', S') that

implements a^* . IR requires that all agents in S' weakly prefer a^* under $u + \tau'$ to b under $u + \tau$, with at least one agent strictly better off. Because b has full coverage, any such strict improvement for some agent must be financed by reducing the utility of other agents at a^* . The key step is to show that this redistribution cannot be supported without violating IR.

Lemma 1 implies that any agent whose utility at a^* is reduced below their base utility must be actively participating in the deviation, i.e., must belong to S' . Hence, any agents outside S' cannot supply the utility needed to fund a strict improvement for members of S' . This forces the total utility at a^* under $u + \tau'$ to exceed the total utility of b for agents in S' , contradicting full coverage, which guarantees that b weakly dominates a^* in total utility. Thus, no IR-feasible transition from (P, τ, S) back to an a^* -implementing state exists. The full proof is presented in Appendix B.1. $\square$

As discussed above, a state (P, τ, S) can be destabilized in two ways. An *outcome-changing* deviation alters reports so that $f(P') \neq f(P)$. An *outcome-preserving* deviation leaves the selected alternative unchanged but allows a coalition to re-allocate transfers so as to increase its members' welfare while maintaining IR and full coverage. Lemma 3 characterizes exactly when such outcome-preserving deviations are feasible.

The key insight is that any profitable redistribution must reduce the utility of some *receiver* agent i with $\tau_i(b) > 0$, where $b = f(P)$. Whether this reduction is feasible depends on the slack in $U_i(b)$ relative to i 's next-best alternatives. If $U_i(b)$ strictly exceeds all other utilities, slack is realized. Otherwise, slack is only potential and must be unlocked by adjusting transfers at tied next-best alternatives. By Lemma 1, any agent whose utility falls below their base utility must participate in the deviation. As a result, slack can be extracted from i only if transfers at i 's next-best alternatives can be relaxed without forcing additional agents into the coalition. Lemma 3 identifies three mutually exclusive scenarios in which this is possible, depending on whether slack is realized, supported by loose donors, or blocked by binding donors.

Definition 5 (Next-best alternative). Given base utilities u , consider a state (P, τ, S) with $f(P) = b$ so that b has full coverage under $u + \tau$. For every $i \in \mathcal{N}$, let

$$NBA_i = \arg \max_{a \neq b} U_i(a)$$

denote the set of next-best alternatives to b .

Definition 6 (Binding NBA alternatives). Given base utilities u , consider a state (P, τ, S) with $f(P) = b$ so that b has full coverage under $u + \tau$. Fix a receiver i , with $\tau_i(b) > 0$, and $\tau_i(x) > 0$ for all $x \in NBA_i$ such that $U_i(b) = U_i(x)$. For any $x \in NBA_i$, let

$$J_x = \{j \in \mathcal{N} : \tau_j(x) < 0\};$$

$$\bar{J}_x = \{j \in J_x : U_j(x) = U_j(b)\}.$$

We denote by

$$\mathcal{B}_i = \{x \in NBA_i : \bar{J}_x = J_x\}$$

the set of *binding* NBA alternatives where all donors are tight.

Lemma 3. Given base utilities u and consensus rule f , consider a state (P, τ, S) that is IR-feasible from the truthful state $(P^*, \emptyset, \emptyset)$. Suppose that $f(P^*) = a^* \notin \arg \max_{a \in \mathcal{A}} SW(a)$, that $f(P) = b \in \arg \max_{a \in \mathcal{A}} SW(a)$, and that b has full coverage under $u + \tau$. There exists an IR-feasible coalitional deviation (P', τ', S') with $f(P') \in \arg \max_{a \in \mathcal{A}} SW(a)$ if and only if $\exists i \in \mathcal{N}$ with $\tau_i(b) > 0$ so that at least one of the following hold:

- P1. (Strict slack at b)
 $U_i(b) > \max_{a \in \mathcal{A} \setminus \{b\}} U_i(a)$;
- P2. (Loose donors at every NBA.)
 $\forall x \in NBA_i$, both $\tau_i(x) > 0$ and $\exists j \in \mathcal{N}$ so that $\tau_j(x) < 0$ and $U_j(x) < U_j(b)$;
- P3. (Potential slack blocked only by binding donors.)
Both:
 - $\forall x \in NBA_i$, $\tau_i(x) > 0$;
 - $\exists j \in \mathcal{N}$ so that $\forall x \in \mathcal{B}_i$, $\tau_j(x) < 0$ and $U_j(x) = U_j(b)$.

Proof Sketch. The proof proceeds in two directions.

Sufficiency. The three cases correspond to progressively weaker ways of extracting transferable slack from a receiver at b , while preserving IR and full coverage constraints throughout the deviation. Fix a state (P, τ, S) with $f(P) = b$ and full coverage. Suppose $\exists i$ with $\tau_i(b) > 0$ satisfying one of (P1)-(P3). In each case, we construct a small redistribution of transfers that lowers $U_i(b)$ by some $\epsilon > 0$ and reallocates this slack to other agents, while preserving both IR for the deviating coalition and full coverage of b .

Under (P1), agent i has *realized* slack: $U_i(b)$ strictly exceeds all other alternatives. A donor at b can therefore withdraw transfers from i without changing i 's ranking, and capture the slack while maintaining IR.

Under (P2), slack is only *potential*: $U_i(b)$ ties with each next-best alternative. However, at every such alternative $x \in NBA_i$ there exists a donor j_x , $\tau_{j_x}(x) < 0$, whose slack is not tight, with $U_{j_x}(x) < U_{j_x}(b)$. These donors can each withdraw a small amount of utility from i at their respective alternatives without violating IR or full coverage of b , collectively freeing slack at b that can be redistributed.

Under (P3), some NBAs are *binding*: all donors at those alternatives are tight. Slack can still be extracted if there exists a single agent who is tight across all binding alternatives. This agent can absorb the necessary adjustment across those alternatives and receive the released slack at b .

In all cases, the deviating coalition excludes i , so the reduction in i 's utility does not violate IR, and the modified transfers preserve full coverage of b . Hence an IR-feasible deviation with $f(P') = b$ exists.

Necessity. Conversely, suppose there exists an IR-feasible deviation (P', τ', S') with $f(P') \in \arg \max_{a \in \mathcal{A}} SW(a)$. Since b has full coverage, any strict gain for agents in S' must be financed by reducing the utility of some non-participating receiver k with $\tau_k > 0$. By Lemma 1, this reduction must preserve full coverage for k at all alternatives.

If k has realized slack at b , then (P1) holds. Otherwise, slack must come from k 's next-best alternatives. If at every such alternative there is a loose donor, then (P2) holds.

If some alternatives are binding, feasibility requires a single agent who can absorb the required adjustments across all binding alternatives, yielding (P3). If none of these conditions holds, no IR-preserving redistribution can reduce $U_k(b)$, contradicting existence of the deviation.

Thus, conditions (P1)–(P3) are exhaustive and characterize exactly when an IR-feasible deviation preserving a welfare-maximizing outcome exists. The full proof is provided in Appendix B.2. $\square$

Our results combine to yield a complete characterization of IR-SNE. The key conclusion is that the truthful state is an IR-SNE if and only if the truthful winner is a social-welfare maximizer. Intuitively, Lemma 2 rules out deviations that revert the outcome to a non-welfare-maximizing alternative once a welfare-maximizing outcome with full coverage is reached. Lemma 3 then characterizes when coalitions can profitably reallocate transfers while preserving the same welfare-maximizing outcome. Condition 3 of the theorem below encodes the slack constraints identified in Lemma 3, ensuring that no such welfare-preserving deviation can exist.

Moreover, no IR-feasible deviation can select an outcome outside $\arg \max_{a \in \mathcal{A}} SW(a)$. Any such deviation would necessarily violate full coverage. In particular, Lemma 5 in Appendix B.1 implies that full coverage pins down utilities across all welfare-maximizing alternatives, ruling out IR-feasible deviations to non-welfare-maximizing outcomes.

Theorem 1 (Characterization of IR-Strong Nash Equilibria). *Fix base utilities u and consensus rule f . A state (P, τ, S) is an IR-SNE if and only if all of the following hold:*

1. **(IR-feasibility.)** (P, τ, S) is IR-feasible from the truthful state $(P^*, \emptyset, \emptyset)$ (or $(P, \tau, S) = (P^*, \emptyset, \emptyset)$).
2. **(Efficiency and coverage.)** $f(P) \in \arg \max_{a \in \mathcal{A}} SW(a)$ and $f(P)$ has full coverage under $u + \tau$.
3. **(No profitable welfare-preserving deviation.)** For every $i \in \mathcal{N}$ so that $\tau_i(f(P)) > 0$, none of the conditions of Lemma 3 are satisfied.

Theorem 1 yields three transparent and easy-to-check conditions for verifying whether a state (P, τ, S) constitutes an IR-SNE. In particular, it separates the problem of equilibrium existence into feasibility, outcome selection, and the absence of profitable welfare-preserving deviations. This structure naturally suggests a constructive approach.

Algorithm 1 exploits this characterization to explicitly construct an IR-SNE by iteratively building transfers that secure full coverage for a welfare-maximizing alternative while eliminating all exploitable slack. The following proposition establishes correctness and efficiency of this construction.

Proposition 3. *Given base utilities u and consensus rule f with default $a^* \in \mathcal{A}$, consider the truthful state $(P^*, \emptyset, \emptyset)$ with $f(P^*) = a^* \notin \arg \max_{a \in \mathcal{A}} SW(a)$. Let $D_t = \{i \in \mathcal{N} : u_i(t) > U_i(t+1)\}$. Algorithm 1 yields a contract scheme τ so that (P, τ, S) is an IR-SNE, where $f(P) \in \arg \max_{a \in \mathcal{A}} SW(a)$ and*

$$S = \bigcup_{t < a^*} D_t \cup \left\{ i \in \mathcal{N} : \arg \max_{a \in \mathcal{A}} U_i(a) \neq \arg \max_{a \in \mathcal{A}} u_i(a) \right\}$$

Algorithm 1: IR-SNE Construction for Consensus

Require: Utilities $(u_i(a))_{i \in \mathcal{N}, a \in \mathcal{A}}$, consensus rule f with $f(P^*) = a^* \notin \arg \max_{a \in \mathcal{A}} SW(a)$
Ensure: Transfer scheme τ with (P, τ, S) is an IR-SNE

- 1: Initialize $\tau_i(a) \leftarrow 0$ for all $i \in \mathcal{N}, a \in \mathcal{A}$
- 2: Order alternatives by social welfare: $1 \geq 2 \geq \dots \geq m$
- 3: **for all** $i \in \mathcal{N}$ **do**
- 4: $\tau_i(a^*) \leftarrow 0; U_i(a^*) \leftarrow u_i(a^*)$
- 5: **end for**
- 6: **for all** $i \in \mathcal{N}, t = a^* + 1, a^* + 2, \dots, m$ **do**
- 7: $U_i(t) \leftarrow \frac{SW(t)}{SW(a^*)} u_i(a^*); \tau_i(t) \leftarrow U_i(t) - u_i(t)$
- 8: **end for**
- 9: **for** $t = a^* - 1, a^* - 2, \dots, 1$ **do**
- 10: Let $R_t = \{i \in \mathcal{N} : u_i(t) < U_i(t+1)\};$
 $D_t = \{i \in \mathcal{N} : u_i(t) > U_i(t+1)\}$
- 11: **for all** $i \in R_t$ **do**
- 12: $\tau_i(t) \leftarrow U_i(t+1) - u_i(t)$
- 13: **end for**
- 14: Distribute total deficit $-\sum_{i \in R_t} \tau_i(t)$ among agents in D_t so that for all $i \in D_t, u_i(t) + \tau_i(t) \geq U_i(t+1)$
- 15: **for all** $i \in \mathcal{N}$ **do**
- 16: $U_i(t) \leftarrow u_i(t) + \tau_i(t)$
- 17: **end for**
- 18: **end for**

is the set of agents that actively participate in the algorithm via donations or strategic voting. It runs in $\mathcal{O}(nm)$ time.

Proof Sketch. Algorithm 1 constructs the transfer scheme τ in a stepwise manner to ensure that the welfare-maximizing alternative b becomes the equilibrium outcome, while the default alternative a^* receives no transfers.

Intuitively, lines 3–5 assign $\tau_i(a^*) = 0$ for all $i \in \mathcal{N}$, so that a^* cannot benefit from additional contributions. Lines 6–8 ensure that a^* weakly dominates every alternative with lower social welfare in $u + \tau$, so they do not pose a risk of deviation and can be ignored for equilibrium considerations. Since a^* is not a social welfare maximizer, lines 9–18 then iterate through alternatives $t \in \{a^* - 1, \dots, 1\}$, partitioning agents into *receivers* R_t (who have smaller utility on t than on $t + 1$) and *donors* D_t (who have strictly greater utility). Receivers are granted transfers to make up their deficit, while donors contribute to satisfy budget balance and maintain full coverage of t .

By construction, any agent i with $\tau_i(b) > 0$ does not satisfy any of the slack conditions (P1)–(P3) from Lemma 3. Combined with Lemma 2, which rules out IR-feasible deviations to a^* , and Lemma 5, which rules out deviations to non-welfare-maximizing alternatives, we conclude that no IR-feasible deviation exists.

A full verification of conditions (P1)–(P3) for all agents is provided in Appendix B.1. $\square$

3.3 Stability Criteria for General AMR Rules

Extending the characterization of IR-SNE beyond the consensus rule is challenging. For many AMR rules, such as Maximin and Ranked Pairs, even deciding whether a fixed coal-

tion can manipulate an outcome is NP-hard [Xia *et al.*, 2009]. Related hardness results extend to several bribery variants when the number of alternatives is unbounded [Faliszewski, 2008; Faliszewski *et al.*, 2009]. Rather than attempting a full characterization under a fixed voting rule, we instead identify a necessary condition for coalitional stability that is independent of the particular AMR rule and can be checked efficiently.

Proposition 1 implies that any IR-SNE outcome must select a SW-maximizing alternative. This leaves a natural question: *even among welfare maximizers, when can a coalition reallocate transfers to induce a profitable deviation?* To address this question, we introduce a notion that captures how much negative transfer can be reassigned without violating individual rationality.

Definition 7 (Reallocatable Amount). Given a state (P, τ, S) and an alternative $\hat{a} \in \mathcal{A}$, the *reallocatable amount* to agent $j \in \mathcal{N}$ from a set of donors $S \subseteq \mathcal{N}$ over $\hat{a}$ is

$$\mathcal{RA}(j, S, \hat{a}) = \sum_{i \in S} \max\{0, -\tau_i(\hat{a})\} + \max\{0, -\tau_j(\hat{a})\}.$$

Intuitively, the reallocatable amount measures how much utility can be redirected toward a target agent from agents who are indifferent between the current outcome and a candidate alternative, together with the agent’s own negative transfers. The following proposition uses this notion to characterize when an IR-feasible deviation exists under *some* AMR rule. Crucially, the voting rule is allowed to be chosen adversarially, isolating a purely economic constraint on stability.

Proposition 4 (IR-Stability via Indifferent Donors). *Fix base utilities u and consider a state (P, τ, S) that is IR-feasible from the truthful state $(P^*, \emptyset, \emptyset)$. Suppose that $f(P^*) = a^* \notin \arg \max_{a \in \mathcal{A}} SW(a)$, that $f(P) = b \in \arg \max_{a \in \mathcal{A}} SW(a)$, and that b has full coverage under $u + \tau$.*

Then there exists an AMR rule f and an IR-feasible state (P', τ', S') with $f(P') = \hat{a} \in \mathcal{A} \setminus \{b\}$ if and only if there exists $j \in \mathcal{N}$ such that

$$\mathcal{RA}(j, S_j(\hat{a}), \hat{a}) > U_j(b) - U_j(\hat{a}),$$

where

$$S_j(\hat{a}) = \{i \in \mathcal{N} \setminus \{j\} : U_i(b) = U_i(\hat{a}), \tau_i(\hat{a}) < 0\}.$$

Proof Sketch. The key idea is that any IR-feasible deviation toward an alternative $\hat{a}$ requires contributions from agents who are indifferent between $\hat{a}$ and the current winner b . The *reallocatable amount* $\mathcal{RA}(j, S_j(\hat{a}), \hat{a})$ measures the maximum utility that can be redirected to a beneficiary j using these indifferent donors and j ’s own negative transfers. If this amount exceeds the difference $U_j(b) - U_j(\hat{a})$, a suitably chosen AMR rule can make $\hat{a}$ the winner, providing sufficiency. Conversely, if no agent satisfies this inequality, any attempt to fund a gain at $\hat{a}$ would require contributions from agents who strictly prefer b , violating IR; this establishes necessity. Thus, the reallocatable-amount condition captures precisely when a coalition can profitably induce an alternative outcome under some AMR rule. The full proof is in Appendix B.2. $\square$

Proposition 4 demonstrates that full coverage alone is insufficient for IR-SNE. In addition to ensuring $U_i(f(P)) \geq$

$U_i(f(P^*))$ for all agents, stability also requires that no agent can be subsidized via reallocations from indifferent donors to overturn the outcome. As an immediate corollary, Lemma 2 follows by setting $\hat{a} = a^*$, since the reallocatable-amount condition fails for all agents.

When the AMR rule f is fixed, passing the reallocatable-amount test is no longer sufficient to guarantee an IR-feasible deviation, reflecting the classical computational hardness of manipulation and bribery problems. For example, swap bribery remains NP-complete for k -approval with $k \geq 3$ [Elkind *et al.*, 2009]. Nevertheless, the reallocatable-amount inequality serves as a polynomial-time filter for candidate deviations, ruling out many alternatives that cannot be sustained under any AMR rule.

4 Discussion

In this paper, we introduced an endogenous bribery mechanism that allows agents to coordinate plans in the presence of coalitional strategic behavior, along with a solution concept that is computationally simple and guaranteed to exist under the consensus rule. At the same time, our results expose several limitations and open problems that must be addressed before such mechanisms can be deployed in practice.

First, multi-agent planning is inherently combinatorial, since each agent may need to take different actions, leading to an exponential number of plans. As a result, computing IR-SNE may be intractable in realistic settings. This challenge motivates future work on simplifying the representation of the problem, exploiting structure in preferences, and developing approximation algorithms that can scale to larger instances.

Second, the welfare consequences of introducing an exchange market warrant further investigation. Although the proposed mechanism explicitly accounts for coalitional strategic behavior, the resulting equilibria may reduce overall efficiency or fairness compared to outcomes where agents act independently or within small cliques of agents [Kearns, 2008]. These concerns are particularly salient when agents have unequal pricing power and when fairness is a key design consideration. Promising directions for future research include extending the framework to online settings, where agents or alternatives arrive and depart over time, as well as exploring auction-based mechanisms for contract formation.

Finally, our analysis adopts Pareto dominance as the basis for individual rationality, implicitly assuming that agents are willing to participate in any coalition that weakly benefits them. This abstraction overlooks the possibility that agents may object to outcomes in which others receive disproportionately large gains, particularly when they have strategic leverage through credible non-participation. For example, consider two agents with default payoffs $(0, 0)$ and two feasible outcomes: $(A) = (1, 1000)$ and $(B) = (2, 100)$. Although both outcomes are Pareto optimal, agent 1 may refuse to participate unless agent 2 agrees to outcome (B) , illustrating how leverage and commitment power can shape equilibrium selection. Modeling such behavior is closely related to the problem of optimal commitment, which is computationally hard and remains an important direction for future work [Littman and Stone, 2001; Conitzer and Sandholm, 2006].

References

- [Amato, 2024] Christopher Amato. An introduction to centralized training for decentralized execution in cooperative multi-agent reinforcement learning. *arXiv preprint arXiv:2409.03052*, 2024.
- [Axelrod and Hamilton, 1981] Robert Axelrod and William D Hamilton. The evolution of cooperation. *science*, 211(4489):1390–1396, 1981.
- [Bachrach *et al.*, 2009] Yoram Bachrach, Edith Elkind, Reshef Meir, Dmitrii Pasechnik, Michael Zuckerman, Jörg Rothe, and Jeffrey S Rosenschein. The cost of stability in coalitional games. In *International Symposium on Algorithmic Game Theory*, pages 122–134. Springer, 2009.
- [Bachrach *et al.*, 2011] Yoram Bachrach, Edith Elkind, and Piotr Faliszewski. Coalitional voting manipulation: A game-theoretic perspective. In *Proceedings of the Twenty-Second International Joint Conference on Artificial Intelligence*, volume 10, pages 978–981, 2011.
- [Baumeister and Rothe, 2016] Dorothea Baumeister and Jörg Rothe. Preference aggregation by voting. In *Economics and Computation: An Introduction to Algorithmic Game Theory, Computational Social Choice, and Fair Division*, pages 197–325. Springer, 2016.
- [Brelsford *et al.*, 2008] Eric Brelsford, Piotr Faliszewski, Edith Hemaspaandra, Henning Schnoor, and Ilka Schnoor. Approximability of manipulating elections. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 8, pages 44–49, 2008.
- [Casella and Macé, 2021] Alessandra Casella and Antonin Macé. Does vote trading improve welfare? *Annual Review of Economics*, 13(1):57–86, 2021.
- [Casella and Palfrey, 2021] Alessandra Casella and Thomas R. Palfrey. Trading votes for votes: A laboratory study. *Games and Economic Behavior*, 125:1–26, 2021.
- [Casella *et al.*, 2012] Alessandra Casella, Aniol Llorente-Saguer, and Thomas R Palfrey. Competitive equilibrium in markets for votes. *Journal of Political Economy*, 120(4):593–658, 2012.
- [Chalkiadakis *et al.*, 2011] Georgios Chalkiadakis, Edith Elkind, and Michael Wooldridge. *Computational aspects of cooperative game theory*. Morgan & Claypool Publishers, 2011.
- [Christodoulou and Koutsoupias, 2005] George Christodoulou and Elias Koutsoupias. The price of anarchy of finite congestion games. In *Proceedings of the Thirty-Seventh Annual ACM Symposium on Theory of Computing*, pages 67–73, 2005.
- [Conitzer and Sandholm, 2006] Vincent Conitzer and Tuomas Sandholm. Computing the optimal strategy to commit to. In *Proceedings of the 7th ACM Conference on Electronic Commerce*, pages 82–90, 2006.
- [Conitzer *et al.*, 2007] Vincent Conitzer, Tuomas Sandholm, and Jérôme Lang. When are elections with few candidates hard to manipulate? *Journal of the ACM*, 54(3):14–es, 2007.
- [Elkind *et al.*, 2009] Edith Elkind, Piotr Faliszewski, and Arkadii Slinko. Swap bribery. In *International Symposium on Algorithmic Game Theory*, pages 299–310. Springer, 2009.
- [Ephrati and Rosenschein, 1997] Eithan Ephrati and Jeffrey S Rosenschein. A heuristic technique for multi-agent planning. *Annals of Mathematics and Artificial Intelligence*, 20(1):13–67, 1997.
- [Faliszewski and Rothe, 2016] Piotr Faliszewski and Jörg Rothe. Control and bribery in voting. In Felix Brandt, Vincent Conitzer, Ulle Endriss, Jérôme Lang, and Ariel D. Procaccia, editors, *Handbook of Computational Social Choice*, pages 146–168. Cambridge University Press, 2016.
- [Faliszewski *et al.*, 2006] Piotr Faliszewski, Edith Hemaspaandra, and Lane A Hemaspaandra. The complexity of bribery in elections. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 6, pages 641–646, 2006.
- [Faliszewski *et al.*, 2007] Piotr Faliszewski, Edith Hemaspaandra, Lane A. Hemaspaandra, and Jörg Rothe. Llull and copeland voting broadly resist bribery and control. In *Proceedings of the Twenty-Second AAAI Conference on Artificial Intelligence*, pages 724–730. AAAI Press, 2007.
- [Faliszewski *et al.*, 2009] Piotr Faliszewski, Edith Hemaspaandra, and Lane A Hemaspaandra. How hard is bribery in elections? *Journal of Artificial Intelligence Research*, 35:485–532, 2009.
- [Faliszewski *et al.*, 2011] Piotr Faliszewski, Edith Hemaspaandra, and Lane A Hemaspaandra. Multimode control attacks on elections. *Journal of Artificial Intelligence Research*, 40:305–351, 2011.
- [Faliszewski *et al.*, 2015] Piotr Faliszewski, Yannick Reisch, Jörg Rothe, and Lena Schend. Complexity of manipulation, bribery, and campaign management in bucklin and fallback voting. *Autonomous Agents and Multi-Agent Systems*, 29(6):1091–1124, 2015.
- [Faliszewski, 2008] Piotr Faliszewski. Nonuniform bribery. In *Proceedings of the 7th International Joint Conference on Autonomous Agents and Multiagent Systems*, pages 1569–1572. International Foundation for Autonomous Agents and Multiagent Systems, 2008.
- [Fershtman *et al.*, 1991] Chaim Fershtman, Kenneth L Judd, and Ehud Kalai. Observable contracts: Strategic delegation and cooperation. *International Economic Review*, pages 551–559, 1991.
- [Geffner *et al.*, 2025] Ivan Geffner, Caspar Oesterheld, and Vincent Conitzer. Maximizing social welfare with side payments. *arXiv preprint arXiv:2508.07147*, 2025.

- [Grandi *et al.*, 2019] Umberto Grandi, Davide Grossi, and Paolo Turrini. Negotiable votes. *Journal of Artificial Intelligence Research*, 64:895–929, 2019.
- [Green and Laffont, 1977] Jerry Green and Jean-Jacques Laffont. Characterization of satisfactory mechanisms for the revelation of preferences for public goods. *Econometrica: Journal of the Econometric Society*, pages 427–438, 1977.
- [Haupt *et al.*, 2024] Andreas Haupt, Phillip Christoffersen, Mehul Damani, and Dylan Hadfield-Menell. Formal contracts mitigate social dilemmas in multi-agent reinforcement learning. *Autonomous Agents and Multi-Agent Systems*, 38(2):51, 2024.
- [Jackson and Wilkie, 2005] Matthew O Jackson and Simon Wilkie. Endogenous games and mechanisms: Side payments among players. *The Review of Economic Studies*, 72(2):543–566, 2005.
- [Kalai and Kalai, 2013] Adam Kalai and Ehud Kalai. Cooperation in strategic games revisited. *The Quarterly Journal of Economics*, 128(2):917–966, 2013.
- [Kearns, 2008] Michael Kearns. Graphical games. In *The New Palgrave Dictionary of Economics*, pages 1–4. Springer, 2008.
- [Kolumbus *et al.*, 2024] Yoav Kolumbus, Joe Halpern, and Éva Tardos. Paying to do better: Games with payments between learning agents. *arXiv preprint arXiv:2405.20880*, 2024.
- [Kraus, 1997] Sarit Kraus. Negotiation and cooperation in multi-agent environments. *Artificial Intelligence*, 94(1-2):79–97, 1997.
- [Lin, 2012] Andrew Peter Lin. *Solving Hard Problems in Election Systems*. PhD thesis, Rochester Institute of Technology, 2012.
- [Littman and Stone, 2001] Michael L Littman and Peter Stone. Implicit negotiation in repeated games. In *International Workshop on Agent Theories, Architectures, and Languages*, pages 393–404. Springer, 2001.
- [Lowe *et al.*, 2017] Ryan Lowe, Yi I Wu, Aviv Tamar, Jean Harb, Pieter Abbeel, and Igor Mordatch. Multi-agent actor-critic for mixed cooperative-competitive environments. In *Advances in Neural Information Processing Systems*, 2017.
- [Meir, 2017] Reshef Meir. Iterative voting. *Trends in Computational Social Choice*, 4:69–86, 2017.
- [Monderer and Tennenholtz, 2009] Dov Monderer and Moshe Tennenholtz. Strong mediated equilibrium. *Artificial Intelligence*, 173(1):180–195, 2009.
- [Nath and Sandholm, 2019] Swaprava Nath and Tuomas Sandholm. Efficiency and budget balance in general quasi-linear domains. *Games and Economic Behavior*, 113:673–693, 2019.
- [Parkes and Xia, 2012] David Parkes and Lirong Xia. A complexity-of-strategic-behavior comparison between schulze’s rule and ranked pairs. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 26, pages 1429–1435, 2012.
- [Sandholm *et al.*, 1999] Tuomas Sandholm, Kate Larson, Martin Andersson, Onn Shehory, and Fernando Tohmé. Coalition structure generation with worst case guarantees. *Artificial Intelligence*, 111(1-2):209–238, 1999.
- [Sodomka *et al.*, 2013] Eric Sodomka, Elizabeth Hilliard, Michael Littman, and Amy Greenwald. Coco-q: Learning in stochastic games with side payments. In *International Conference on Machine Learning*, pages 1471–1479. PMLR, 2013.
- [Thomson, 1994] William Thomson. Cooperative models of bargaining. In Robert J. Aumann and Sergiu Hart, editors, *Handbook of Game Theory with Economic Applications*, volume 2, pages 1237–1284. Elsevier, 1994.
- [Turrini, 2013] Paolo Turrini. Endogenous boolean games. In *Proceedings of the Twenty-Third International Joint Conference on Artificial Intelligence*, pages 390–396, 2013.
- [Xia *et al.*, 2009] Lirong Xia, Michael Zuckerman, Ariel D. Procaccia, Vincent Conitzer, and Jeffrey S. Rosenschein. Complexity of unweighted coalitional manipulation under some common voting rules. In *Proceedings of the 21st International Joint Conference on Artificial Intelligence*, pages 348–353, 2009.
- [Xia, 2012] Lirong Xia. Computing the margin of victory for various voting rules. In *Proceedings of the 13th ACM Conference on Electronic Commerce*, pages 982–999, 2012.
- [Yang *et al.*, 2020] Jiachen Yang, Ang Li, Mehrdad Farajtabar, Peter Sunehag, Edward Hughes, and Hongyuan Zha. Learning to incentivize other learning agents. In *Advances in Neural Information Processing Systems*, volume 33, pages 15208–15219, 2020.