

Repeated Nash implementation

CLAUDIO MEZZETTI

School of Economics, University of Queensland

LUDOVIC RENO

School of Economics and Finance, Queen Mary University of London

We study the repeated implementation of social choice functions in environments with complete information and changing preferences. We define *dynamic monotonicity*, a natural but nontrivial dynamic extension of Maskin monotonicity, and show that it is necessary and almost sufficient for repeated Nash implementation, regardless of whether the horizon is finite or infinite and whether the discount factor is “large” or “small.”

KEYWORDS. Dynamic monotonicity, Nash implementation, Maskin monotonicity, repeated implementation, repeated games.

JEL CLASSIFICATION. C72, D71.

1. INTRODUCTION

Many economic and social interactions are repeated: the same buyers and sellers often trade with one another multiple times, teams of contractors regularly work for the same procurement agencies, and voters repeatedly elect representatives, to name just a few. The central theme of this paper is the design of institutions, or contractual arrangements, that generate “socially desirable” outcomes in settings where agents repeatedly interact and preferences change over time.

To illustrate the type of economic problems this paper addresses, consider for example the situation in which a buyer and a seller interact more than once. Are there contractual arrangements that (in all equilibria) allow the seller to extract all the surplus from trade? As another example, consider the case in which two (or more) agents may work on a number of tasks that are profitable to a principal. Can we design arrangements that (again, in all equilibria) induce the agents to work on the most profitable tasks at each point in time, even if it is costly to them? In all these problems, an essential difficulty is the multiplicity of equilibria, including “undesirable” equilibria, that repeated interactions make possible to sustain. The aim of the paper is to characterize the social outcomes that are implementable; that is, those outcomes for which there exist contractual arrangements that only yield equilibria consistent with them.

Claudio Mezzetti: c.mezzetti@uq.edu.au

Ludovic Renou: lrenou.econ@gmail.com

We would like to thank an associate editor and anonymous referees for their insightful comments and suggestions. Mezzetti’s work was funded in part by Australian Research Council Grant DP120102697.

Copyright © 2017 The Authors. Theoretical Economics. The Econometric Society. Licensed under the Creative Commons Attribution-NonCommercial License 3.0. Available at <http://econtheory.org>.

DOI: 10.3982/TE1988

More formally, we study the problem of repeated, full implementation of social choice functions in environments with complete information and a changing state of the world. A social choice function is repeatedly implementable in Nash equilibrium if there exists a sequence of (possibly history-dependent) mechanisms such that for any period, for any profile of preferences at that period, the set of equilibrium outcomes corresponds to the social choice function at that profile of preferences.

Full implementation in a static environment (i.e., with a single period) has been extensively studied.¹ The seminal contribution is [Maskin \(1999\)](#), which states that Maskin monotonicity is necessary and almost sufficient for full implementation. In this paper, we provide a condition, called *dynamic monotonicity*, and show that it is necessary and almost sufficient for repeated Nash implementation, regardless of whether the horizon is finite or infinite and whether the discount factor is “large” or “small.”

Dynamic monotonicity is a natural but nontrivial dynamic extension of Maskin monotonicity. It reduces to Maskin monotonicity in single-period settings, but is weaker in all other finitely repeated implementation problems. Thus, perhaps surprisingly, finitely repeated implementation is “easier” to achieve than single-shot implementation. For example, while full-surplus extraction by a seller cannot be implemented in a static problem, it can be if there are at least two periods in which the buyer and the seller interact (see [Example 1 in Section 3](#)).

We also show that in infinitely repeated problems with patient enough agents, dynamic monotonicity implies that the social choice function is weakly efficient from the agents’ point of view. However, no efficiency condition is necessary in infinitely repeated problems with impatient enough agents and in all finitely repeated problems. For example, collusion among agents in a team can be deterred in all finite horizon problems and in infinite horizon problems with impatient enough agents (see [Example 2 in Section 3](#)).

In a repeated implementation problem, the designer’s choice of a mechanism in each period may depend on the agents’ actions and mechanisms in all previous periods; agents need not be playing the same stage game in each period. Intuitively, contractual arrangements may be used to compensate an agent when he deviates before period t from a collusive strategy profile that would induce socially undesirable outcomes from period t onward. This possibility of inducing preemptive deviations from future collusion facilitates implementation and is the reason why finitely repeated implementation is easier than static implementation. Indeed, it is only when the horizon is infinite and the discount factor is close to 1 that the gain from a future collusive agreement dominates any preemptive punishment and only outcomes that are efficient for the agents can be implemented. This insight is at the heart of [Lee and Sabourian’s \(2011\)](#) work on infinitely repeated implementation problems (to be discussed shortly).

Unlike the literature on dynamic mechanism design, which has recently seen a flurry of papers (e.g., see the survey by [Bergemann and Said 2011](#)), the literature on full implementation in dynamic environments is in its infancy. Two papers have studied repeated setting where, unlike in this paper, the state of the world does not change over time.

¹See [Jackson \(2001\)](#), [Maskin and Sjöström \(2002\)](#), and [Serrano \(2004\)](#) for recent surveys on implementation theory.

Kalai and Ledyard (1998) study infinitely repeated implementation in dominant strategies; they show that every social choice function can be repeatedly implemented starting from some (possibly distant) point in the future. Chambers (2004) studies virtual repeated Nash implementation in continuous time.

In an important recent paper, Lee and Sabourian (2011) consider environments in which, like in our paper, the state of the world changes over time.² Unlike us, they focus on infinitely repeated settings with patient agents; that is, agents with a discount factor arbitrarily close to 1. Their main result is that weak efficiency of the social choice function relative to any other function with an equal or smaller range is necessary for infinitely repeated implementation. Under some mild additional assumptions on the environment, they also show that if the discount factor is larger than $1/2$, then strict efficiency in the range is sufficient for infinitely repeated implementation from period two onward (but the designer may fail to implement the correct outcome in the first period).

Maskin monotonicity and weak efficiency in the range are very different conditions, and thus it is perhaps a puzzle that the first is necessary and almost sufficient in the static case and the second is necessary and almost sufficient in the polar case of infinite interactions with patient enough agents. In this paper, we solve this puzzle by introducing the condition of dynamic monotonicity and showing that it is necessary and almost sufficient in all repeated implementation problems, including the so-far unexplored, but clearly empirically important, case of a finite number of interactions and the case of infinitely repeated interactions with general discount factors. In the static case, dynamic monotonicity is equivalent to Maskin monotonicity. In infinitely repeated problems with an arbitrarily high enough discount factor, dynamic monotonicity is essentially equivalent to weak efficiency in the range. As we illustrate in Examples 1 and 2, neither Maskin monotonicity nor an efficiency condition are necessary for repeated implementation in general.

The paper is organized as follows. Section 2 defines the problem of repeated implementation. Section 3 presents two examples motivating our investigation. Section 4 introduces the condition of dynamic monotonicity. Section 5 presents the main results of the paper. Section 6 provides some extensions of our results and Section 7 concludes. All proofs are given in the Appendix.

2. DEFINITIONS

Single-shot implementation. A static or *single-shot* implementation problem $\mathcal{P}$ is a tuple $\langle \mathcal{I}, X, \Theta, (u_i)_{i \in \mathcal{I}} \rangle$, where $\mathcal{I} = \{1, \dots, I\}$ is a set of I agents, X is the set of alternatives—a compact subset of a finite dimensional Euclidean space, Θ is a finite set of states of the world, and for each agent $i \in \mathcal{I}$, $u_i : X \times \Theta \rightarrow \mathbb{R}$ is a state-dependent continuous utility function. Let $L_i(x, \theta) = \{y \in X : u_i(x, \theta) \geq u_i(y, \theta)\}$ be agent i 's lower contour set of x at state θ . A social choice function (henceforth, scf) $f : \Theta \rightarrow X$ associates with each state of the world θ the alternative $f(\theta) \in X$.

²See also Renou and Tomala (2015) and Lee and Sabourian (2013) for the problem of approximate implementation in environments with incomplete information.

A static mechanism G is a pair $((M_i^G)_{i \in \mathcal{I}}, g)$ where M_i^G is the set of messages of agent i and $g : \times_{i \in \mathcal{I}} M_i^G \rightarrow X$ is the allocation rule. Let $M^G = \times_{j \in \mathcal{I}} M_j^G$ and $M_{-i}^G = \times_{j \in \mathcal{I} \setminus \{i\}} M_j^G$, where m and m_{-i} are generic elements. The mechanism $\langle M^G, g \rangle$ and the state θ induce the strategic-form game $G(\theta) = \langle \mathcal{I}, (u_i(g(\cdot), \theta), M_i^G)_{i \in \mathcal{I}} \rangle$. Let $\text{NE}^G(\theta) \subseteq X$ be the set of (pure) Nash equilibrium outcomes of the game $G(\theta)$. The social choice function f is single-shot implementable in Nash equilibrium if there exists a static mechanism G such that $\text{NE}^G(\theta) = \{f(\theta)\}$ for all $\theta \in \Theta$.

A necessary and almost sufficient condition for static Nash implementation is *Maskin monotonicity*. In Definition 1, we present two equivalent, slightly unusual, formulations of Maskin monotonicity, as they foreshadow and will help understanding our definition of dynamic monotonicity. Call any map $\pi : \Theta \rightarrow \Theta$ a (static) *deception* and let Π^1 be the set of static deceptions. The interpretation is that when the state is θ , agents act as if the state were $\pi(\theta)$ instead.

DEFINITION 1. A social choice function f is Maskin monotonic when it satisfies (M^A) or, equivalently, (M^B) .

(M^A) For all $\pi \in \Pi^1$, for all $\theta \in \Theta$,

$$[\forall i \in \mathcal{I}, L_i(f(\pi(\theta)), \pi(\theta)) \subseteq L_i(f(\pi(\theta)), \theta)] \Rightarrow [f(\pi(\theta)) = f(\theta)].$$

(M^B) For all $\pi \in \Pi^1$, for all $\theta \in \Theta$,

$$\begin{aligned} & [f(\pi(\theta)) \neq f(\theta)] \\ & \Rightarrow [\exists (i \in \mathcal{I}, x \in X) : u_i(f(\pi(\theta)), \pi(\theta)) - u_i(x, \pi(\theta)) \\ & \qquad \qquad \qquad \geq 0 > u_i(f(\pi(\theta)), \theta) - u_i(x, \theta)]. \end{aligned}$$

The intuition for the necessity of Maskin monotonicity is simple. Suppose that f is implementable and let π be a deception. At state $\pi(\theta)$, there must exist an equilibrium m^* that implements $f(\pi(\theta))$. However, if $f(\pi(\theta)) \neq f(\theta)$, m^* should not be an equilibrium at state θ , so that at least one agent must have a profitable deviation; that is, he must have a unilateral deviation from m^* that induces an alternative x strictly preferred to $f(\pi(\theta))$ at state θ . And since m^* is an equilibrium at state $\pi(\theta)$, the deviation cannot be profitable at $\pi(\theta)$; that is, $f(\pi(\theta))$ is preferred to x at state $\pi(\theta)$. Condition (M^B) precisely captures this intuition.

Repeated implementation. A *repeated implementation problem*, denoted $\mathcal{P}^T$, represents the T -time repetition of the implementation problem $\mathcal{P}$; T can be finite or infinite. At the beginning of each period $t \in \mathcal{T} = \{1, \dots, T\}$, the state of the world is drawn from Θ with probability mass function p , with $p(\theta) > 0$ for all $\theta \in \Theta$. In each period, the realized state is commonly observed by all agents, but not the designer.

Let $(x(t, \theta))_{t \in \mathcal{T}, \theta \in \Theta}$ be a sequence of alternatives, where $x(t, \theta)$ is the alternative implemented in state θ at period t . An agent's expected payoff over sequences of alternatives is given by the discounted criterion; that is, there exists $\delta \in (0, 1)$ such that the

payoff of agent i from $(x(t, \theta))_{t \in \mathcal{T}, \theta \in \Theta}$ is given by³

$$U_i((x(t, \theta))_{t \in \mathcal{T}, \theta \in \Theta}) = \frac{1 - \delta}{1 - \delta^T} \sum_{t \in \mathcal{T}} \sum_{\theta \in \Theta} \delta^{t-1} u_i(x(t, \theta), \theta) p(\theta).$$

The aim of the designer is to repeatedly implement a social choice function f . A dynamic mechanism regime specifies a mechanism in each period t , contingent on the profile of mechanisms offered and messages played up to period t (excluding period t). A designer history h_D^t is a sequence of mechanisms and corresponding messages $(G_1, m_1, \dots, G_\tau, m_\tau, \dots, G_{t-1}, m_{t-1})$ such that G_τ is the mechanism adopted at period τ and $m_\tau \in M^{G_\tau}$ is the corresponding message profile, for all $\tau < t$. The set of all possible histories observed by the designer at period t is denoted $\mathcal{H}_D^t$. The set of initial histories $\mathcal{H}_D^1$ is the singleton $\{\emptyset\}$ and the set of all possible designer histories is $\mathcal{H}_D = \bigcup_{t=1}^T \mathcal{H}_D^t$.

A *dynamic mechanism regime*, or *regime* for short, specifies a lottery over static mechanisms as a function of the designer history. We write $r(G; h_D^t)$ for the probability that mechanism G is chosen after history h_D^t .⁴

We assume perfect monitoring.⁵ At the beginning of period t , each agent knows the entire profile of mechanisms chosen up to period $t - 1$, the entire profile of messages sent up to period $t - 1$, the entire profile of states of the world realized up to period $t - 1$, and the period t 's mechanism selected as well as the realized state of the world for period t . Write $\theta^t = (\theta_1, \dots, \theta_{t-1})$ for a profile of realized states of the world up to period $t - 1$. A history for agent i is thus $h^t = (h_D^t, \theta^t)$. Let $\mathcal{H}^t$ be the set of all possible t -period agent histories and let $\mathcal{H} = \bigcup_{t=1}^T \mathcal{H}^t$ be the set of all such histories. The only possible initial history is the empty set: $\mathcal{H}^1 = \{\emptyset\}$.

A pure strategy s_i for agent i specifies a message in each period t as a function of the history h^t , the mechanism G_t currently selected, and the current state θ_t ; that is, $s_i(h^t, G_t, \theta_t) \in M_i^{G_t}$ for all (h^t, G_t, θ_t) . Let $s = (s_1, \dots, s_I)$ be a strategy profile. The strategy profile s , the random draw of a state in each period, and the regime r generate a random sequence of histories h^t .

Given a regime r , we write $q(h^t; s)$ for the probability that history h^t occurs when the strategy profile is s . Throughout, we slightly abuse notation and write $r(G_t; h^t)$ for $r(G_t; h_D^t)$ for any $h^t = (h_D^t, \theta^t)$. The expected payoff of agent i when the profile of strategies is s is

$$U_i(s) = \frac{1 - \delta}{1 - \delta^T} \sum_{t \in \mathcal{T}} \sum_{h^t \in \mathcal{H}^t} \sum_{G_t \in \mathcal{G}} \sum_{\theta_t \in \Theta} \delta^{t-1} u_i(g(s(h^t, G_t, \theta_t), \theta_t) q(h^t; s) r(G_t; h^t) p(\theta_t).$$

³When computing payoffs starting from any period t , we use the normalizing factor $(1 - \delta)/(1 - \delta^{T-t+1})$, so that the discounted payoff from t is measured on the same scale as the single-shot payoff.

⁴We assume that, for each h_D^t , $r(\cdot; h_D^t)$ has finite support.

⁵In other words, we assume that the designer truthfully and publicly reveals all his information (i.e., messages received, alternative implemented, and mechanism selected) at each period. In a more general model, the communication policy would also be part of the design problem, i.e., the designer would also choose how much to reveal to the agents in each period. Clearly, this can only enlarge the set of implementable social choice functions.

A profile of pure strategies $s^* = (s_i^*, s_{-i}^*)$ is a pure Nash equilibrium of the dynamic game induced by regime r if $U_i(s_i^*) \geq U_i(s_i, s_{-i}^*)$ for all strategies s_i , for all agents $i \in \mathcal{I}$ (where s_{-i}^* denotes the strategy profile of agent i 's opponents).

DEFINITION 2. A social choice function f is repeatedly implementable if there exists a dynamic mechanism regime r such that (i) there exists a Nash equilibrium s^* of the dynamic game induced by r and (ii) for each Nash equilibrium s induced by r , we have $g(s(h^t, G_t, \theta_t)) = f(\theta_t)$ for all $\theta_t \in \Theta$, for all (h^t, G_t) such that $q(h^t; s) > 0$, and $r(G_t; h^t) > 0$, for all $t \in \mathcal{T}$.

Intuitively, a social choice function is repeatedly implementable if we can construct a dynamic mechanism whose unique equilibrium outcome is $f(\theta)$ in all periods where the state is θ . As is customary in the literature, [Definition 2](#) does not rule out mixed strategy equilibria with outcome realizations different from $f(\theta)$. In [Section 6](#) we will show that it is possible to rule out such undesirable mixed strategy equilibria.

We end this section with three notions of efficiency of an scf. The expected payoff of agent i when f is repeatedly implemented is $v_i^f = \sum_{\theta \in \Theta} u_i(f(\theta), \theta) p(\theta)$. Let $\mathcal{F}(f) = \{f' : \Theta \rightarrow X : f'(\Theta) \subseteq f(\Theta)\}$ be the set of social choice functions with a range (weakly) smaller than f , let $V(f) = \{(v_i)_{i \in \mathcal{I}} : v_i = v_i^{f'} \text{ for all } i \in \mathcal{I}, \text{ for some } f' \in \mathcal{F}(f)\}$ be the associated (expected) payoff profiles, and let $\text{co}(V(f))$ be the convex hull of $V(f)$.

The social choice function f is *weakly efficient in the range* if there does not exist a payoff profile $(v_i)_{i \in \mathcal{I}} \in \text{co}(V(f))$ such that $v_i > v_i^f$ for all $i \in \mathcal{I}$; f is *efficient in the range* if there does not exist a payoff profile $(v_i)_{i \in \mathcal{I}} \in \text{co}(V(f))$ such that $v_i \geq v_i^f$ for all $i \in \mathcal{I}$ and $v_i > v_i^f$ for some $i \in \mathcal{I}$; f is *strictly efficient in the range* if it is efficient in the range and there exists no $f' \in \mathcal{F}(f)$, $f' \neq f$, such that $v_i^{f'} = v_i^f$ for all $i \in \mathcal{I}$.⁶

3. TWO EXAMPLES

This section illustrates repeated implementation with the help of two simple examples.

EXAMPLE 1 (Trading a Good). This is a multiperiod variation of the leading example of [Aghion et al. \(2012\)](#).

There are two periods, $t = 1, 2$, a buyer B , and a seller S . In each period, the seller has a good for sale; the quality θ of the good is independently drawn in each period and equally likely to be $\theta_L = 10$ or $\theta_H = 14$. The buyer and the seller have a common discount factor δ and observe the good's quality at the beginning of each period.

As in [Aghion et al.](#), payments to and from a third party are allowed. Hence, the set of outcomes X is the set of triplets (z, p_B, p_S) with $z \in \{0, 1\}$ representing whether the good is traded ($z = 1$) or not ($z = 0$), $p_B \in P$ representing the price paid by the buyer, and $p_S \in P$ representing the price paid to the seller, where P is a (arbitrarily large) closed interval in $\mathbb{R}$. For any outcome (z, p_B, p_S) , the (per-period) buyer's utility is $u(z\theta - p_B)$

⁶Efficiency and strict efficiency in the range were first defined by [Lee and Sabourian \(2011\)](#).

		Seller			
		θ_L	$N\theta_L$	$N\theta_H$	θ_H
Buyer	θ_L	(1, 10, 10) $\hookrightarrow G_2$	(1, 10, 10) $\hookrightarrow (1, 11, 11)$	(0, 1, 1) $\hookrightarrow (0, 0, 0)$	(0, 1, 1) $\hookrightarrow (0, 0, 0)$
	$N\theta_L$	(1, 14, 10) $\hookrightarrow (0, -y, 0)$	(1, 14, 14) $\hookrightarrow (1, 14, 14)$	(1, 0, 0) $\hookrightarrow (1, 0, 0)$	(0, 1, 1) $\hookrightarrow (0, 0, 0)$
	$N\theta_H$	(0, 1, 1) $\hookrightarrow (0, 0, 0)$	(1, 0, 0) $\hookrightarrow (1, 0, 0)$	(1, 14, 14) $\hookrightarrow (1, 14, 14)$	(1, 14, 10) $\hookrightarrow (0, 0, 0)$
	θ_H	(0, 1, 1) $\hookrightarrow (0, 0, 0)$	(0, 1, 1) $\hookrightarrow (0, 0, 0)$	(1, 14, 14) $\hookrightarrow (1, 11, 11)$	(1, 14, 14) $\hookrightarrow G_2$

TABLE 1. The first-period allocation and the transition $\hookrightarrow$ to the second-period allocation or game played in Example 1.

when the good quality is θ , with $u(0) = 0$ and u a strictly increasing, strictly concave function. The seller’s utility is p_S .

We want to implement the efficient allocation prescribing that in each period the good is traded and the buyer pays the seller the true quality, $p_B = p_S = \theta$; that is, the scf we want to implement is $f(\theta_L) = (1, 10, 10)$ and $f(\theta_H) = (1, 14, 14)$.⁷ Since f is not Maskin monotonic, it cannot be implemented in Nash equilibrium in a static setting.⁸

We now present a simple dynamic mechanism that repeatedly implements f in Nash equilibrium. In the first period, the buyer and the seller report a message in $\{\theta_L, N\theta_L, N\theta_H, \theta_H\}$. We interpret the report θ_k as stating that “the quality is θ_k .” The reports $N\theta_k$ are objections that lead to different first-period allocations and second-period mechanisms than announcing either θ_H or θ_L . In the second period, the buyer and the seller have the opportunity to make an additional report in $\{\theta_L, \theta_H\}$ if and only if they have reported the same quality in the first period. In all other cases, the second-period allocation is chosen without requiring buyer and seller to make reports. Table 1 gives the allocation rule in the first period along with the regime.

Table 1 has 16 cells, one for each possible report profile in the first period; the row (resp., column) report is the buyer (resp., seller) report. Each cell has two elements. The top element gives the first-period allocation, while the bottom element (indicated with the symbol $\hookrightarrow$) gives the transition to the second-period mechanism. For instance, if the buyer reports θ_L and the seller reports $N\theta_L$, the first-period allocation is $(1, 10, 10)$, while the second-period mechanism implements $(1, 11, 11)$ and requires no second-period reports. When the buyer and the seller report the same quality in the first period, the second-period mechanism is G_2 , given in Table 2:⁹

⁷Note that if $u'(0) = 1$, then this allocation also maximizes total surplus.

⁸Formally, we have that $L_B(f(\theta_L), \theta_L) = \{(z, p_B, p_S) : u(0) \geq u(z\theta_L - p_B)\} \subseteq \{(z, p_B, p_S) : u(4) \geq u(z\theta_H - p_B)\} = L_B(f(\theta_L), \theta_H)$, while $L_S(f(\theta_L), \theta_L) = L_S(f(\theta_L), \theta_H)$. Since $f(\theta_L) \neq f(\theta_H)$, we have a violation of Maskin monotonicity.

⁹On the equilibrium path, mechanism G_2 guarantees that trade takes place in the second period and the expected price is 12 for both the buyer and the seller.

	θ_L	θ_H
θ_L	(1, 10, 10)	(0, 0, 0)
θ_H	(0, 0, 0)	(1, 14, 14)

TABLE 2. The second-period mechanism G_2 in Example 1.

We claim that whenever y is chosen so that $-u(-4) > \delta u(y) > u(4)$, the unique pure strategy equilibrium implements the efficient allocation in both periods.¹⁰ This is verified in the Appendix, which presents the two reduced strategic-form games that are obtained by conditioning on the first-period quality.¹¹

A notable feature of our mechanism is that it provides at least one agent with the incentive to deviate early (at $t = 1$) from future (at $t = 2$) coordination on undesirable equilibria (coordinating on announcing θ_L when the good’s quality is θ_H). It is precisely the ability to provide such incentives in a dynamic setting that allows the repeated Nash implementation of social choice functions, like the one in this example, that are not implementable in a static setting. $\diamond$

EXAMPLE 2 (Task assignment). In each of a possibly infinite number of periods, a principal needs to assign two agents (experts), 1 and 2, to one of two tasks, A and B . There are two states of the world, $\theta \in \{\theta_A, \theta_B\}$. The agents know the state of the world, but not the principal. In state θ_A (resp., θ_B), task A (resp., B) yields the principal a benefit v greater than the cost to undertake it, while the other task yields zero benefit and cost. An allocation is a quadruplet (a_1, a_2, w_1, w_2) , with $a_i \in \{A, B\}$ the assignment of agent $i \in \{1, 2\}$ and $w_i \geq 0$ his wage. When the state is θ , the assignment is (a_i, a_{-i}) , and the wage is w_i , agent i ’s payoff is $w_i - c_i(a_i, a_{-i}, \theta)$, where $c_i(a_i, a_{-i}, \theta)$ is agent i ’s cost of executing task a_i when the other agent is assigned to task a_{-i} , at state θ . There are complementarities: the more agents work on a task, the less costly it is: $c_i(a_i, a_{-i}, \theta) = 1$ if $(a_i, a_{-i}, \theta) = (A, A, \theta_A) = (B, B, \theta_B)$, $c_i(a_i, a_{-i}, \theta) = 3$ if $(a_i, a_{-i}, \theta) = (A, B, \theta_A) = (B, A, \theta_B)$, and the cost is zero otherwise; in addition, v is sufficiently large, e.g., $v > 4$, so that it is profitable for the principal to induce the agents to work on the right task.

The principal wants to maximize his ex post profit in each period, subject to giving the agents at least their per-period outside option payoff, which we normalize to zero. This corresponds to the scfs $f(\theta_A) = (A, A, 1, 1)$ and $f(\theta_B) = (B, B, 1, 1)$. Note that f maximizes social surplus in each period and state.

The scf f is Maskin monotonic, but it is not efficient relative to social choice functions having (weakly) smaller ranges. For instance, the functions $f^*(\theta_A) = (B, B, 1, 1)$ and $f^*(\theta_B) = (A, A, 1, 1)$, with agents being paid to work on the unprofitable task, give a strictly higher expected utility to both agents than f . Thus, if the agents are sufficiently patient, then f cannot be repeatedly implemented in infinite horizon problems (Theorem 1, Lee and Sabourian 2011).

¹⁰The existence of y follows from observing that $u(4) + u(-4) < 0$, since u is strictly concave and $u(0) = 0$.
¹¹As we argue in Section 6, undesirable mixed strategy equilibria could also be ruled out, at the cost of introducing a more complicated mechanism.

		Agent 2	
		θ_A	θ_B
Agent 1	θ_A	$(A, A, 1, 1)$	$(B, A, 2, 2)$
	θ_B	$(A, B, 2, 2)$	$(B, B, 1, 1)$

TABLE 3. The static mechanism in Example 2.

	θ_A		θ_B	
	Agent 1	Agent 2	Agent 1	Agent 2
$(A, A, 1, 1)$	0	0	1	1
$(B, B, 1, 1)$	1	1	0	0
$(B, A, 2, 2)$	2	−1	−1	2
$(A, B, 2, 2)$	−1	2	2	−1

TABLE 4. Agents’ payoffs in Example 2.

At the end of Section 5 we will show that f is infinitely repeatedly implementable if the discount factor is not too large. We now argue that the f is repeatedly implementable in any finite horizon problem. Consider the static mechanism where each agent has two messages, θ_A and θ_B , and the allocation rule is represented as in Table 3; Table 4 displays the payoffs to each agent of each alternative in each state.

At state θ_A , the mechanism induces a prisoner’s dilemma, with (θ_A, θ_A) as the unique Nash equilibrium and equilibrium outcome $(A, A, 1, 1)$. Similarly, at state θ_B , the mechanism induces a prisoner’s dilemma, with (θ_B, θ_B) as the unique Nash equilibrium and equilibrium outcome $(B, B, 1, 1)$. So f is implementable when $T = 1$. More fundamentally, at states θ_A and θ_B , the unique equilibrium payoff coincides with the min-max payoff. Consequently, repeated play of the stage game equilibrium is the only Nash equilibrium of the finitely repeated game (e.g., see Benoît and Krishna, 1987, and González-Díaz 2006), and by selecting the mechanism regime that uses the static mechanism in each round, f can be finitely repeatedly implemented in Nash equilibrium, regardless of the number of periods.

This shows that there is an important difference between what can be implemented in finitely repeated problems and what can be implemented in infinitely repeated problems with an arbitrarily large discount factor, as studied by Lee and Sabourian (2011). $\diamond$

4. DYNAMIC MONOTONICITY

Consider any period t and any sequence $(u_i^\tau)_{\tau \geq t}$ of payoffs from period t onward. We can write agent i ’s discounted payoff at period t as

$$\frac{1 - \delta}{1 - \delta^{T-t+1}} \left(u_i^t + \delta \sum_{\tau=t+1}^T \delta^{\tau-t-1} u_i^\tau \right) = (1 - \beta_{t,T}) u_i^t + \beta_{t,T} v_i(t),$$

where $v_i(t)$ is the (normalized) discounted continuation payoff and $\beta_{t,T}$ is the (normalized) discount factor at period t : that is,

$$v_i(t) = \frac{1 - \delta}{1 - \delta^{T-t}} \sum_{\tau=t+1}^T \delta^{\tau-t-1} u_i^\tau \quad \text{and} \quad \beta_{t,T} = \frac{\delta - \delta^{T-t+1}}{1 - \delta^{T-t+1}}.$$

When the horizon is infinite, i.e., $T = \infty$, we have $\beta_{t,\infty} = \delta$. The lowest and highest expected payoff agent i can obtain are

$$\underline{v}_i = \sum_{\theta \in \Theta} \min_{x \in X} u_i(x, \theta) p(\theta), \quad \bar{v}_i = \sum_{\theta \in \Theta} \max_{x \in X} u_i(x, \theta) p(\theta).$$

For each $t \in \mathcal{T} \setminus \{T\}$, let $V_i(t)$ be the closed interval $[\underline{v}_i, \bar{v}_i]$ with the convention that $V_i(T) = \{0\}$ if $T < \infty$. The set $V_i(t)$ corresponds to the set of feasible agent i 's (normalized) continuation payoffs at period t . Denote by $v_i^f(t)$ the (normalized) expected discounted payoff of agent i when f is implemented from period $t + 1$ onward. Thus, $v_i^f(t) = v_i^f = \sum_{\theta \in \Theta} u_i(f(\theta), \theta) p(\theta)$ if $t < T$ and $v_i^f(T) = 0$ if $T < \infty$.

We now generalize the important concept of *deception* to the dynamic setting. At each period t , a deception specifies a state $\hat{\theta}_t$ as a function of the realized state θ_t and the history of realized states up to period t , θ^t . Formally, a deception π is a sequence of maps $(\pi_t : \Theta^t \times \Theta \rightarrow \Theta)_{t=1}^T$. Intuitively, suppose that each agent is asked to directly report a state at each period (as in a direct mechanism). A deception then corresponds to a situation where the agents coordinate their reports to $\hat{\theta}_t = \pi_t(\theta^t, \theta_t)$ at period t , when the current state is θ_t and the profile of realized states is θ^t .¹² (If reports are not coordinated, the designer detects a lie and can punish the agents.) Of course, the mechanism does not have to be direct. Nonetheless, the concept of a deception remains important: agents can play at period t and realized states θ^t as if the current state is $\pi_t(\theta^t, \theta_t)$ and not θ_t . A special deception is π^* , given by $\pi_t^*(\theta^t, \theta_t) = \theta_t$ for all (θ^t, θ_t) , for all t . This corresponds to truth-telling. Let Π^T be the set of deceptions.

We define the (normalized) expected discounted continuation payoff of agent i from following the deception π after state history (θ^t, θ_t) recursively as

$$\begin{aligned} v_i^{f,\pi}(\theta^t, \theta_t) = & \sum_{\theta_{t+1} \in \Theta} \left((1 - \beta_{t+1,T}) u_i(f(\pi_{t+1}((\theta^t, \theta_t), \theta_{t+1})), \theta_{t+1}) \right. \\ & \left. + \beta_{t+1,T} v_i^{f,\pi}((\theta^t, \theta_t), \theta_{t+1}) \right) p(\theta_{t+1}). \end{aligned}$$

This is agent i 's discounted continuation payoff if, in all periods $\tau > t$, the designer uses the social choice function f at the reported state $\pi_\tau(\theta^\tau, \theta_\tau)$ to determine the period τ alternative. Note that the discounted continuation payoff $v_i^{f,\pi^*}(\theta^t, \theta_t)$ from the truth-telling deception π^* is equal to $v_i^f(t)$, regardless of (θ^t, θ_t) .

¹²Note that π and the history of realized states θ^t determine a unique history of reported states.

For any history of realized states θ^t and deception π , we define the dynamic lower contour set of x at θ_t as

$$\begin{aligned} L_{i,\theta^t}^{f_\pi}(x, \theta_t) &= \{(y, v_i(t)) \in X \times V_i(t) : (1 - \beta_{t,T})u_i(y, \theta_t) + \beta_{t,T}v_i(t) \\ &\leq (1 - \beta_{t,T})u_i(x, \theta_t) + \beta_{t,T}v_i^{f_\pi}(\theta^t, \theta_t)\}. \end{aligned}$$

Dynamic lower contour sets are defined in the space of alternatives and continuation payoffs. Intuitively, for any deception π and history of states θ^t , the dynamic lower contour set at θ_t is composed of all the pairs of alternatives and continuation payoffs that give agent i a smaller expected discounted payoff than when x is implemented at state θ_t in period t and agents continue to follow the deception π from period $t + 1$ onward. Note that $L_{i,\theta^t}^{f_\pi}(x, \theta_t)$ does not depend on θ^t , since the truth-telling deception π^* does not. With a slight abuse of notation, we therefore write $L_{i,t}^{f_\pi}(x, \theta_t)$ for $L_{i,\theta^t}^{f_\pi}(x, \theta_t)$.

We are now ready to present two equivalent definitions of dynamic monotonicity, the dynamic generalization of Maskin monotonicity.

DEFINITION 3 (Dynamic monotonicity). A social choice function f is dynamic monotonic if it satisfies (DM^A) or, equivalently, (DM^B).

(DM^A) For all $\pi \in \Pi^T$, for all $\theta^T \in \Theta^T$,

$$\begin{aligned} [\forall(i \in \mathcal{I}, t \in \mathcal{T}), L_{i,t}^{f_\pi}(f(\pi_t(\theta^t, \theta_t)), \pi_t(\theta^t, \theta_t)) \subseteq L_{i,t}^{f_\pi}(f(\pi_t(\theta^t, \theta_t)), \theta_t)] \\ \Rightarrow [\forall(t \in \mathcal{T}), f(\pi_t(\theta^t, \theta_t)) = f(\theta_t)]. \end{aligned}$$

(DM^B) For all $\pi \in \Pi^T$, for all $\theta^T \in \Theta^T$,

$$\begin{aligned} [\exists(t' \in \mathcal{T}) : f(\pi_{t'}(\theta^{t'}, \theta_{t'})) \neq f(\theta_{t'})] \\ \Rightarrow [\exists(i \in \mathcal{I}, t \in \mathcal{T}, x \in X, v_i \in V_i(t)) : \\ (1 - \beta_{t,T})[u_i(f(\pi_t(\theta^t, \theta_t)), \pi_t(\theta^t, \theta_t)) - u_i(x, \pi_t(\theta^t, \theta_t))] \\ + \beta_{t,T}[v_i^f(t) - v_i] \geq 0, \\ 0 > (1 - \beta_{t,T})[u_i(f(\pi_t(\theta^t, \theta_t)), \theta_t) - u_i(x, \theta_t)] \\ + \beta_{t,T}[v_i^{f_\pi}(\theta^t, \theta_t) - v_i]]. \end{aligned}$$

Intuitively, dynamic monotonicity says that if agents coordinate on a deception that induces an undesirable alternative at some period t' (for some profile of realized states), then at least one agent must have a profitable deviation starting at some time t . Since the problem is dynamic, the profitable deviation does not have to start at t' ; it could start before or after; t need not equal t' . For instance, in [Example 1](#), the seller has a profitable deviation at the first period from the second-period coordination on trading the high quality good at the low price.

It is worth noting that we can restrict attention to deceptions that weakly dominate truth-telling in checking for dynamic monotonicity, i.e., to deceptions π such that $v_i^{f_\pi}(\theta^t, \theta_t) \geq v_i^f$ for all i , for all (θ^t, θ_t) , for all t .

A few additional observations are worth making. First, for $T = 1$, dynamic monotonicity reduces to Maskin monotonicity. Second, observe that when $T = \infty$, $\beta_{t,T} = \delta$ for all t , and the dynamic lower contour sets do not vary with t . Consequently, when checking for dynamic monotonicity, it is sufficient to consider $t = 1$. Third, an easy-to-check sufficient condition for dynamic monotonicity is as follows. For each agent i , define $v_i^{\max} = \max_{\pi: \Theta \rightarrow \Theta} \sum_{\theta} u_i(f(\pi(\theta)), \theta) p(\theta)$ as the highest payoff that agent i can obtain if all agents coordinate on the most favorable static deception π for agent i ($v_i^{\max}$ is also the highest payoff that agent i can obtain by maximizing over all dynamic deceptions). Suppose that $f(\theta) \neq f(\theta^*)$. Using (DM^B), a sufficient condition for dynamic monotonicity is that for all deceptions such that $\pi_{t'}(\theta', \theta^*) = \theta$ for some $\theta' \in \Theta'$ and $t' \in \mathcal{T}$, there exist $t \in \mathcal{T}$, $i \in \mathcal{I}$, $x \in X$, and $v_i(t) \in V_i(t)$ that satisfy

$$(1 - \beta_{t,T})u_i(f(\theta), \theta) + \beta_{t,T}v_i^f \geq (1 - \beta_{t,T})u_i(x, \theta) + \beta_{t,T}v_i(t)$$

and

$$(1 - \beta_{t,T})u_i(f(\theta), \theta^*) + \beta_{t,T}v_i^{\max} < (1 - \beta_{t,T})u_i(x, \theta^*) + \beta_{t,T}v_i(t).$$

The example illustrating [Remark 5](#) shows that this condition is easy to check.

We end this section with a series of remarks. The message we want to convey is that dynamic monotonicity is the “general” condition for repeated Nash implementation. It reduces to Maskin monotonicity when there is a single period and essentially corresponds to [Lee and Sabourian’s \(2011\)](#) efficiency in the range when there are an infinite number of periods and a discount factor close to 1.

The first remark gives another easy-to-check sufficient condition for the dynamic monotonicity of a social choice function. The second remark states that, in finite horizon problems, dynamic monotonicity is weaker than Maskin monotonicity. The converse is false; [Example 1](#) demonstrates that dynamic monotonicity is strictly weaker than Maskin monotonicity.

REMARK 1. If the social choice function f is strictly efficient in the range and $(v_i^f)_{i \in \mathcal{I}}$ is an extreme point of $\text{co}(V(f))$, then f is dynamic monotonic whenever $T \geq 2$.

REMARK 2. Suppose $T < \infty$. If f is Maskin monotonic, then it is dynamic monotonic.

REMARK 3. Suppose $T = \infty$. There exists $\delta_H \in (0, 1)$ such that for all $\delta \in (\delta_H, 1)$, if f is dynamic monotonic, then it is weakly efficient in the range.

REMARK 4. Suppose $T = \infty$. If f is Maskin monotonic and efficient in the range, then it is dynamic monotonic.

REMARK 5. There are social choice functions, which are neither efficient nor Maskin monotonic, and yet are dynamically monotonic.

As a demonstration of [Remark 5](#), suppose that there are two agents, two periods, no discounting (i.e., $\delta = 1$), two equiprobable states of the world θ and θ' , and five alternatives a, b, c, d, e . Let the payoffs be as in [Table 5](#).

	θ	θ'
a	3, 3	1, 7
b	6, 0	3, 3
c	10, 10	10, 10
d	-10, -10	0, 0
e	0, 0	-10, -10

TABLE 5. Agents' payoffs in the example illustrating Remark 5.

The social choice functions are $f(\theta) = a$ and $f(\theta') = b$, and the associated payoff profile is $(v_1^f, v_2^f) = (3, 3)$. It is not Maskin monotonic since $L_i(f(\theta'), \theta') \subseteq L_i(f(\theta'), \theta)$ for all i , and yet $f(\theta') \neq f(\theta)$. It is also not efficient in the range since if players coordinate on θ (resp., θ') when the state is θ' (resp., θ), then they each obtain a payoff of 7/2. Yet, f is dynamic monotonic. To see this, remember that $v_i^{\max}$ is the highest payoff that agent i can obtain if all agents coordinate on the most favorable deception for agent i , and note that $v_1^{\max} = 9/2$, while $v_2^{\max} = 5$. It is immediate to check that the pair $(d, 10)$ satisfies

$$u_1(f(\theta), \theta) + v_1^f = 3 + 3 \geq -10 + 10 = u_1(d, \theta) + v_1$$

$$\max(u_1(f(\theta), \theta'), u_1(f(\theta'), \theta')) + v_1^{\max} = 3 + 9/2 < 0 + 10 = u_1(d, \theta') + v_1.$$

Similarly, the pair $(e, 10)$ satisfies

$$u_2(f(\theta'), \theta') + v_2^f = 3 + 3 \geq -10 + 10 = u_2(e, \theta') + v_2$$

$$\max(u_2(f(\theta'), \theta), u_2(f(\theta), \theta)) + v_2^{\max} = 3 + 5 < 0 + 10 = u_2(e, \theta) + v_2.$$

We have the necessary preference reversals in the first period and, therefore, the social choice function is dynamic monotonic.

The final remark states that in finitely repeated settings the set of social choice functions that are dynamic monotonic is weakly increasing in T .

REMARK 6. Suppose $T < \infty$ and f is dynamic monotonic over T periods. Then f is also dynamic monotonic over $T + 1$ periods.¹³

5. MAIN RESULTS

This section presents our main results, stating that dynamic monotonicity is necessary and almost sufficient for repeated Nash implementation. We begin with necessity.

THEOREM 1 (Necessity). *If the social choice function f is repeatedly implementable, then it is dynamic monotonic.*

The intuition for Theorem 1 is simple and analogous to the intuition for the necessity of Maskin monotonicity in static implementation problems. If the social choice

¹³We thank an anonymous referee for asking us to verify whether this claim holds.

function f is implementable, there must exist a mechanism and an equilibrium such that $f(\theta_t)$ is implemented at period t and state θ_t , and the continuation payoff to any agent i is $v_i^f(t)$, for any $t \in \mathcal{T}$. Moreover, for any realized profile of states θ^t , all deviations at period t and state θ_t must give to agent i an alternative x and a continuation payoff v_i in $L_{i,\theta^t}^f(f(\theta_t), \theta_t)$. Consider a deception π and a “collusive” equilibrium in which agents follow the deception (on the equilibrium path) and revert to the original equilibrium after unilateral deviations. In particular, agents pretend that the state is $\pi_t(\theta^t, \theta_t^*) = \theta_t$ when the realized state at period t is θ_t^* and the history of realized states up to period t is θ^t . As a result, $f(\theta_t) = f(\pi_t(\theta^t, \theta_t^*))$ is implemented at t in state θ_t^* , and the expected payoff of agent i is $(1 - \beta_{t,T})u_i(f(\theta_t), \theta_t^*) + \beta_{t,T}v_i^{\pi}(\theta^t, \theta_t^*)$. If $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,\theta^t}^{\pi}(f(\theta_t), \theta_t^*)$, then agent i has no profitable deviation from the collusive equilibrium. For otherwise, he would have had a profitable deviation at state θ_t from the original equilibrium. Hence, for f to be implemented, it must be that $f(\theta_t^*) = f(\theta_t)$; that is, f must be dynamic monotonic.

We now consider sufficient conditions. As in static implementation problems, we distinguish between the case of two and more than two agents. We need to introduce some additional definitions.

For each $Y \subseteq X$, define $\max_i^\theta Y = \{x \in Y : u_i(x, \theta) \geq u_i(y, \theta) \text{ for all } y \in Y\}$ as agent i ’s maximal set in Y at state θ . A social choice function f satisfies *no-veto power* if, for all $\theta \in \Theta$, $x \in \max_i^\theta X$ for all $i \in \mathcal{I}^*$ with $|\mathcal{I}^*| \geq I - 1$ implies $f(\theta) = x$. Maskin monotonicity and no-veto power are sufficient for static Nash implementation when there are at least three agents. A similar results holds in the repeated setting once we replace Maskin monotonicity with dynamic monotonicity.

THEOREM 2 (Sufficiency $I \geq 3$). *Let $I \geq 3$. If the social choice function f is dynamic monotonic and satisfies no-veto power, then it is repeatedly implementable.*

The proof is constructive. The main building block of our construction is the static mechanism G^* , a close relative to Maskin’s (1999) canonical mechanism. The mechanism G^* requires the agents to report a state, an alternative, a continuation payoff, and an integer. At period t , “unanimous” reports $(\theta_t, f(\theta_t), v_i^f(t), 0)$ result in the realization of $f(\theta_t)$ and in the adoption of G^* in the next period. A unilateral deviation from unanimity by agent j at t , $(\theta_{j,t}, x_{j,t}, v_{j,t}, n_{j,t})$, results in the realization of $x_{j,t}$ at t and in the continuation payoff v_j^j thereafter, if $(x_{j,t}, v_{j,t})$ is in agent j ’s dynamic contour set $L_{j,t}^f(f(\theta_t), \theta_t)$ (where θ_t is the common state report of all agents but agent j). Alternatively, the deviation results in the realization of $f(\theta_t)$ at period t and in the continuation payoff $v_j^f(t)$ thereafter. To guarantee that agent j obtains $v_{j,t}$ (or $v_j^f(t)$) in the future, the regime appropriately randomizes between adopting a mechanism where agent j is dictatorial (i.e., chooses the alternative), which would guarantee he receives $\bar{v}_j$, and a punishment mechanism where agent j would get less than $v_{j,t}$ (or $v_j^f(t)$). Any other report profile at t leads to the agent reporting the highest integer at t being dictatorial at t and in all future periods. Notice that the mechanism G^* is equivalent to Maskin’s canonical mechanism when $T = 1$, and indeed guarantees the implementation of f for very similar arguments as in Maskin (1999). As the canonical Maskin mechanism with $T = 1$,

our mechanism regime does not rule out undesirable mixed strategy equilibria. As we discuss in [Section 6](#), under a mild additional assumption we can eliminate them.¹⁴

The dynamic mechanism regime we construct only uses stage mechanisms that are deterministic functions of the agents' messages, but permits random transitions between these mechanisms. Without making further assumptions, it seems impossible to prove [Theorem 2](#) without the help of stochastic transitions or, alternatively, stochastic stage mechanisms.¹⁵ Yet, in environments with transfers and quasi-linear preferences, there is no need for stochastic transitions; we can always adjust the transfers to guarantee that the agent obtains the appropriate continuation payoff.

As [Maskin's \(1999\)](#) theorem for static Nash implementation, [Theorem 2](#) requires no-veto power. We can weaken the no-veto power requirement. For instance, [Theorem 2](#) remains valid if we replace no-veto power with [Assumption A](#), stated below, which is closely related to the conditions $\mu(\text{ii})$ and $\mu(\text{iii})$ of Moore and Repullo.¹⁶ We first need some additional notation. Let $\varphi_t : \{t + 1, \dots, T\} \times \Theta \rightarrow X$ be a time-dependent social choice function and write $v_i^{\varphi_t}$ the continuation payoff of implementing φ_t from period $t + 1$ onward, that is,

$$v_i^{\varphi_t} := \frac{1 - \delta}{1 - \delta^{T-t}} \sum_{\tau=t+1}^T \delta^{\tau-t-1} u_i(\varphi_t(\tau, \theta), \theta) p(\theta).$$

For any $v_i \in V_i(t)$, define $\lambda(v_i) = (v_i - \underline{v}_i) / (\bar{v}_i - \underline{v}_i)$. We are now ready to state [Assumption A](#).

ASSUMPTION A. A social choice function f satisfies [Assumption A](#) if the following statements hold:

- (i) For all $(x, v_i(t)) \in L_{i,t}(f(\theta), \theta)$ with $i \in \mathcal{I}$, $\theta \in \Theta$ and $t \in \mathcal{T}$ and for all pairs $(\varphi_t, \bar{\varphi}_t)$ such that
 - (a1) either $\lambda(v_i(t)) = 0$ or $\varphi_t(\tau, \theta) \in \bigcap_j \max_j^\theta X$ for all $\theta \in \Theta$, for all $\tau > t$ ¹⁷
 - (a2) $\bar{\varphi}_t(\tau, \theta) \in \bigcap_{j \neq i} \max_j^\theta X$ for all $\theta \in \Theta$, for all $\tau > t$
 - (b) $x \in \bigcap_{j \neq i} \max_j^{\theta^*} X$
 - (c) $\beta_{t,T} u_i(x, \theta^*) + (1 - \beta_{t,T})[\lambda(v_i(t)) v_i^{\varphi_t} + (1 - \lambda(v_i(t))) \bar{v}_i^{\bar{\varphi}_t}] \geq \beta_{t,T} u_i(y, \theta^*) + (1 - \beta_{t,T}) v_i$ for all $(y, v_i) \in L_{i,t}(f(\theta), \theta)$,

we have that $x = f(\theta^*)$, and $\varphi_t(\tau, \cdot) = \bar{\varphi}_t(\tau, \cdot) = f$ for all $\tau > t$.

- (ii) For all x such that $x \in \bigcap_j \max_j^{\theta^*} X$, we have that $x = f(\theta^*)$.

¹⁴See [Mezzetti and Renou \(2012\)](#) for an alternative definition of static implementation in mixed Nash equilibrium.

¹⁵[Azacis and Vida \(2015\)](#) use random mechanisms and random transitions in their analysis of infinitely repeated implementation problems.

¹⁶We prove this and the following claim in footnotes [22](#) and [23](#).

¹⁷We thank Helmut Azacis and Peter Vida for pointing out the need to add $\lambda(v_i(t)) = 0$ as a special case.

Condition (i) is similar to condition $\mu(ii)$ of Moore and Repullo. It states that if x maximizes the payoff of all agents but agent i at state θ^* , if φ_i maximizes the continuation payoff of all agents while $\bar{\varphi}_i$ maximize the continuation payoff of all agents but agent i , and if the pair $(x, \lambda(v_i(t))v_i^{\varphi_i} + (1 - \lambda(v_i(t)))v_i^{\bar{\varphi}_i})$ is maximal in the dynamic lower contour set $L_{i,t}(f(\theta), \theta)$ at state θ^* , then not only alternative x must coincide with $f(\theta^*)$ at state θ^* , but also $\varphi_i(\tau, \cdot)$ and $\bar{\varphi}_i(\tau, \cdot)$ must coincide with f for all $\tau > t$. Note that condition (i) is weaker than no-veto power and is almost identical to condition $\mu(ii)$ at period T , when $T < \infty$. Condition (ii) is a unanimity condition.

We now consider the two-agent case. As shown by Dutta and Sen (1991) and Moore and Repullo (1988), for the static case with two agents, *self-selection* is a necessary condition for Nash implementation.¹⁸ Our sufficiency result for two agents requires a strengthening of self-selection.¹⁹

ASSUMPTION B. There exists an alternative w such that $u_i(w, \theta) < u_i(f(\theta'), \theta)$ for all $(\theta', \theta) \in \Theta \times \Theta$, for all $i \in \{1, 2\}$.

Assumption B requires that there exists a bad outcome (relative to f) for both agents. For instance, in pure exchange economies with strictly monotone preferences, the zero consumption bundle is a bad outcome relative to any social choice function that gives positive consumption to each consumer in at least one state of the world. Other examples satisfying **Assumption B** include environments with transferable utilities, like our two examples in Section 3. We have the following theorem.

THEOREM 3 (Sufficiency $I = 2$). *Let $I = 2$. Suppose Assumptions A and B hold. If a social choice function f is dynamic monotonic, then it is repeatedly implementable.*

We now briefly return to Examples 1 and 2.

EXAMPLE 1 (revisited). The set $V(f)$ of expected (ex ante) payoff vectors that the two parties would obtain with an scf whose range is a subset of $\{(1, 10, 10), (1, 14, 14)\}$, the range of f , is $\{(u(4)/2, 10), (0, 12), ((u(-4) + u(4))/2, 12), (u(-4)/2, 14)\}$. Thus f , which yields expected payoffs $(v_B^f, v_S^f) = (0, 12)$, is strictly efficient and an extreme point in the convex hull of $V(f)$. By Remark 1, f is dynamic monotonic. Since Assumptions A and B hold, f is repeatedly implementable in Nash equilibrium irrespective of the discount factor, as long as there are at least two periods. $\diamond$

EXAMPLE 2 (revisited). Consider an infinitely repeated setting. To show under which condition f is dynamic monotonic when $T = \infty$, we can use the sufficient condition provided after Definition 3. Observe that $v_i^f = 0$ for all $i \in \mathcal{I}$ and that the best possible collusive deception is $\pi_t(\theta^t, \theta_A) = \theta_B$ and $\pi_t(\theta^t, \theta_B) = \theta_A$ for all θ^t , for all $t \in \mathcal{T}$. Under such a deception, $v_i^{f^\pi}(\theta^t, \theta) = v_i^{f^\pi} = 1$ for all $i \in \mathcal{I}$. (This corresponds to $v_i^{\max}$.) Given the

¹⁸In Proposition 1 in the Appendix, we show that a weaker condition, *dynamic self-selection*, is necessary for repeated Nash implementation.

¹⁹*Self-selection:* Let $I = 2$. There exists $x(\theta_2, \theta_1) \in L_1(f(\theta_2), \theta_2) \cap L_2(f(\theta_1), \theta_1)$ for all pairs (θ_2, θ_1) .

symmetry of the setup, we only need to consider the pairs (θ_A, θ_B) with $\pi_t(\theta^t, \theta_B) = \theta_A$. Since $f(\theta_B) \neq f(\theta_A)$, dynamic monotonicity (DM^B) requires that there exist $i \in \mathcal{I}$, $x \in X$, and $v_i \in V_i(t)$ such that

$$\begin{aligned} (1 - \delta)[u_i(f(\theta_A), \theta_A) - u_i(x, \theta_A)] + \delta[0 - v_i] &\geq 0 \\ 0 &> (1 - \delta)[u_i(f(\theta_A), \theta_B) - u_i(x, \theta_B)] + \delta[1 - v_i]. \end{aligned}$$

This is equivalent to

$$-(1 - \delta)u_i(x, \theta_A) \geq \delta v_i > 1 - (1 - \delta)u_i(x, \theta_B). \quad (1)$$

By symmetry, we may take i to be any agent, say agent 1. The only alternatives x that may satisfy (1) for agent 1 assign agent 1 to task A and agent 2 to task B . Letting $x = (A, B, w_1, w_2)$, (1) becomes $(1 - \delta)(3 - w_1) \geq \delta v_i > 1 - (1 - \delta)w_1$, which holds if and only if $\delta < 2/3$. This shows that f satisfies dynamic monotonicity if the discount factor is less than $2/3$. Thus, dynamic monotonicity does not imply weak efficiency in infinite horizon problems when the discount factor is not too large. Since the setting of the example satisfies Assumptions A and B, with an infinite time horizon, f can be repeatedly implemented, and collusion among the agents avoided, as long as $\delta < 2/3$. $\diamond$

6. DISCUSSION

This section discusses some important aspect of our analysis.

Mixed strategies. The proof of Theorem 2 does not rule out undesirable mixed strategy equilibria. We now show that the theorem extends to mixed strategies under the mild additional assumption of *no indifference*, which states that no agent is totally indifferent between all alternatives at all states.

We say that a scf f is repeatedly implementable in mixed Nash equilibrium if it is repeatedly implementable in Nash equilibrium and, in addition, there are no mixed strategy Nash equilibria that yield in some period t an outcome $y \notin f(\theta)$ with positive probability, when the state is θ .

THEOREM 4. *Let $I \geq 3$. Assume no indifference holds. If the social choice function f is dynamic monotonic and satisfies no-veto power, then it is repeatedly implementable in mixed Nash equilibrium.*

Two obstacles must be overcome when dealing with mixing by agents. First, the best message for an agent to send depends on the messages sent by the other agents, but the agent has no certainty over such messages when the other agents mix. For instance, announcing a large integer so as to become a dictator entails the risk of being the odd man out when others play unanimously. Second, we need to consider distributions over deceptions so as to account for mixed strategies, i.e., distributions over pure strategies. In the proof, we overcome these difficulties by introducing random stage mechanisms. This guarantees that mixing only occurs in the last period in all equilibria (if there is a last period). Moreover, the last-period mechanism is a version of the mechanism in

Maskin and Sjöström (2002), which allows agents to propose alternatives contingent on the state report of their opponents. This guarantees that no undesirable equilibria exist.

Subgame perfection. The solution concept adopted in this paper is Nash equilibrium. All our results extend straightforwardly to subgame perfection. First, it is easy to check that the Nash equilibrium s^E constructed in the proof of Theorem 2 (and Theorem 3) is subgame perfect. Since there are no undesirable Nash equilibria, hence no undesirable subgame-perfect Nash equilibria, this implies that dynamic monotonicity together with no-veto power (or Assumption A) are sufficient for subgame-perfect implementation. Dynamic monotonicity is also necessary as long as the mechanism adopted in each period is a static mechanism. To see this, suppose that f is repeatedly implementable in subgame-perfect Nash equilibrium, and let s be an implementing equilibrium. Assume that there exists a deception π such that for all $t \in \mathcal{T}$, for all $\theta^t \in \Theta^t$, for all pairs (θ_t, θ_t^*) with $\pi_t(\theta^t, \theta_t^*) = \theta_t$, we have $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,\theta_t^*}^{f_\pi}(f(\theta_t), \theta_t^*)$ for all $i \in \mathcal{I}$. As in the proof of Theorem 1, we can construct a Nash equilibrium s' that implements $f(\pi_t(\theta^t, \cdot))$ at all periods t and at all profiles θ^t of realized states up to period t . Moreover, off the equilibrium path, s' agrees with s , so that s' is also a subgame-perfect equilibrium and hence f must be dynamic monotonic.²⁰

Time-dependent social choice functions. We have assumed that the designer wants to implement the same social choice function f in each period. A more general objective would be to implement a sequence $(f_t)_{t \in \mathcal{T}}$ of social choice functions. It is straightforward to modify the definitions of continuation payoffs, dynamic lower contour sets, and dynamic monotonicity to account for time-dependent social choice functions. With these modifications, dynamic monotonicity remains necessary and almost sufficient for repeated Nash implementation.

Social choice correspondences. The analysis extends to the implementation of social choice correspondences. Let $F : \Theta \rightarrow 2^X \setminus \{\emptyset\}$ be a social choice correspondence; denote by $\mathbb{F}$ the set of all possible social choice functions that are selections of F . A social choice correspondence is implementable if there exists a dynamic mechanism such that for every selection $f \in \mathbb{F}$, there exists a Nash equilibrium that repeatedly implements f , and every Nash equilibrium repeatedly implements a selection $f \in \mathbb{F}$. A social choice correspondence F is dynamic monotonic when it satisfies the following criterion:

(DM_C^A) For all $f \in \mathbb{F}$, for all $\pi \in \Pi^T$, for all $\theta^T \in \Theta^T$,

$$\begin{aligned} & [\forall (i \in \mathcal{I}, t \in \mathcal{T}), L_{i,t}^f(f(\pi_t(\theta^t, \theta_t)), \pi_t(\theta^t, \theta_t)) \subseteq L_{i,\theta_t}^{f_\pi}(f(\pi_t(\theta^t, \theta_t)), \theta_t)] \\ & \Rightarrow [\exists f^* \in \mathbb{F} : \forall t \in \mathcal{T}, f(\pi_t(\theta^t, \theta_t)) = f^*(\theta_t)]. \end{aligned}$$

Note that the concept of dynamic monotonicity (for correspondences) is equivalent to Maskin monotonicity (for correspondences) in static implementation problems, and clearly equivalent to Definition 3 when F is single-valued. To see the necessity of

²⁰It is important to stress that the restriction to static mechanisms within a period rules out the mechanisms used by Moore and Repullo (1988) and Abreu and Sen (1990) to show that, in single-shot environments, subgame-perfect implementation is substantially more permissive than Nash implementation.

the modified condition of dynamic monotonicity, suppose that F is repeatedly implementable and assume that there exist a selection $f \in \mathbb{F}$, a deception π such that for all $t \in \mathcal{T}$, for all $\theta^t \in \Theta^t$, for all pairs (θ_t, θ_t^*) with $\pi_t(\theta^t, \theta_t^*) = \theta_t$, we have $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,\theta^t}^{f\pi}(f(\theta_t), \theta_t^*)$ for all $i \in \mathcal{I}$. As in the proof of [Theorem 1](#), we can construct an equilibrium that implements $f(\pi_t(\theta^t, \cdot))$ at all periods t and at all profiles θ^t of realized states up to period t . Consequently, there must exist $f^* \in \mathbb{F}$ such that $f(\pi_t(\theta^t, \cdot)) = f^*$ for all $\theta^t \in \Theta^t$, for all $t \in \mathcal{T}$, i.e., F must be dynamic monotonic. To show sufficiency, we need to augment the dynamic mechanism regime in the proof of [Theorem 2](#) with an initial stage (period $t = 0$) in which all agents announce a selection $f \in \mathbb{F}$. If all agents announce the same selection $f \in \mathbb{F}$ at period $t = 0$, then our dynamic mechanism regime takes effect from $t = 1$ with f the social choice function adopted in the canonical mechanism G_t^* . If not all agents make the same announcement at $t = 0$, then our dynamic mechanism regime takes effect from $t = 1$ with an arbitrary $f^* \in \mathbb{F}$ as the social choice function adopted in G_t^* .

7. CONCLUSIONS

Our main contribution is to introduce the condition of dynamic monotonicity, a natural but nontrivial dynamic extension of Maskin monotonicity, and to show, in [Theorems 1–4](#), that dynamic monotonicity is necessary and almost sufficient for repeated Nash implementation of social choice functions, regardless of whether the horizon is finite or infinite and whether the discount factor is large or small.²¹

Many economic applications of implementation theory, for example most of the contracting literature (e.g., see [Aghion et al., 2012](#) or [Maskin and Tirole 1999](#)), focus on static problems. One of the main insights of our paper is that the (finitely) repeated implementation of desirable social choice functions is easier than static implementation, as last-period, or late periods, planned deviations from truth-telling can be avoided by rewarding defection in early periods. For instance, we can implement full surplus extraction by a seller as long as there are at least two periods, while full surplus extraction is not implementable in static problems (see [Example 1](#)).

APPENDIX

This appendix contains the proofs of all our results and the reduced strategic-form games associated with [Example 1](#).

EXAMPLE 1 (The strategic-form games). Conditional on a realized first-period quality, the buyer and the seller have 64 strategies each. An agent is active at the initial history as well as at the histories (θ_H, θ_H) and (θ_L, θ_L) . At the initial history, the agent has four actions. At histories (θ_H, θ_H) and (θ_L, θ_L) , an agent has two actions for each realization of the second-period quality. All strategies where an agent plays $N\theta_L$ in the first period

²¹Indeed, [Theorem 1](#) also remains true if we adopt a different criterion than the discounting criterion to evaluate streams of payoff, e.g., the overtaking criterion or the limit of the means criterion (naturally, with a modification in the definition of dynamic monotonicity to account for these changes).

are payoff equivalent (there are 16 strategies of that form), and similarly, for all strategies where an agent plays $N\theta_H$ in the first period. We write $N\theta_L$ and $N\theta_H$ for those strategies. If the first-period reports do not match, then the game essentially ends. Thus, all strategies where an agent reports θ_L at the initial history, reports θ_L at the history (θ_L, θ_L) conditional on second-period quality θ_L , and reports θ_L at the history (θ_L, θ_L) conditional on second-period quality θ_H are payoff-equivalent. We write $\theta_L\theta_L\theta_L$ for those strategies. Similarly, for all other strategies. For instance, $\theta_H\theta_H\theta_L$ represents all strategies where an agent reports θ_H at the initial history, reports θ_H at the history (θ_H, θ_H) conditional on second-period quality θ_L , and reports θ_L at the history (θ_H, θ_H) conditional on second-period quality θ_H . Each reduced strategic-form game has therefore 10 “strategies.” Tables 6 and 7 represent the two reduced strategic-form games associated with each first-period quality θ_L and θ_H . The buyer is the row player, while the seller is the column player. In each cell, the top payoff is the buyer’s payoff, while the bottom payoff is the seller’s payoff. $\diamond$

Throughout the proofs, we use the following observation. For any deception $\tilde{\pi} \in \Pi^T$ and state history $\tilde{\theta}^T \in \Theta^T$ such that for all $i \in \mathcal{I}$ and $t \in \mathcal{T}$, $L_{i,t}^f(f(\tilde{\pi}_t(\tilde{\theta}^t, \tilde{\theta}_t)), \tilde{\pi}_t(\tilde{\theta}^t, \tilde{\theta}_t)) \subseteq L_{i,\tilde{\theta}^t}^{\tilde{\pi}}(f(\tilde{\pi}_t(\tilde{\theta}^t, \tilde{\theta}_t)), \tilde{\theta}_t)$, there exists a deception $\pi \in \Pi^T$ such that for all $i \in \mathcal{I}$, $t \in \mathcal{T}$ and, importantly, for all $(\theta^t, \theta_t) \in \Theta^t \times \Theta$, $L_{i,t}^f(f(\pi_t(\theta^t, \theta_t)), \pi_t(\theta^t, \theta_t)) \subseteq L_{i,\theta^t}^{\pi}(f(\pi_t(\theta^t, \theta_t)), \theta_t)$. The deception π agrees with $\tilde{\pi}$ at $\tilde{\theta}^T$ and with π^* at all other state histories. Thus, if f is dynamic monotonic, then $f(\pi_t(\theta^t, \theta_t)) = f(\theta_t)$ for all $t \in \mathcal{T}$ and $(\theta^t, \theta_t) \in \Theta^t \times \Theta$. As the converse is also true, we have an equivalent formulation of dynamic monotonicity.

PROOF OF REMARK 1. Note that since f is strictly efficient in the range, for each $v \in \text{co}(V(f))$ such that $v \neq v^f = (v_i^f)_{i \in \mathcal{I}}$, there exists $i^* \in \mathcal{I}$ such that $v_{i^*} < v_{i^*}^f$. Moreover, since $v = \sum_{f' \in \mathcal{F}(f)} \alpha^{f'} v^{f'}$ with $\sum_{f' \in \mathcal{F}(f)} \alpha^{f'} = 1$ and $\alpha^{f'} \geq 0$ for all $f' \in \mathcal{F}(f)$, it follows from strict efficiency of f and the fact that v^f is an extreme point of $\text{co}(V(f))$ that $\alpha^f = 1$ whenever $v = v^f$, i.e., v corresponds to the implementation of f . Consequently, for any deception π such that $\pi_t(\theta^t, \theta_t^*) = \theta_t \neq \theta_t^*$, $v_{i^*}^{\pi} \in \text{co}(V(f))$ and $v_{i^*}^{\pi} \neq v_{i^*}^f$. Therefore, for some i^* we have $v_{i^*}^{\pi} < v_{i^*}^f$ and hence $(f(\theta_t), v_{i^*}^f) \in L_{i^*,t}^f(f(\theta_t), \theta_t)$, but $(f(\theta_t), v_{i^*}^f) \notin L_{i^*,\theta^t}^{\pi}(f(\theta_t), \theta_t^*)$. $\square$

PROOF OF REMARK 2. Suppose that f is Maskin monotonic and assume that there exists a deception π such that for all $t \in \mathcal{T}$, for all $\theta^t \in \Theta^t$, for all pairs (θ_t, θ_t^*) with $\pi_t(\theta^t, \theta_t^*) = \theta_t$, we have $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,\theta^t}^{\pi}(f(\theta_t), \theta_t^*)$ for all $i \in \mathcal{I}$. We need to show that $f(\theta_t^*) = f(\theta_t)$ for all $\theta^t \in \Theta^t$, for all $t \in \mathcal{T}$. The argument is by induction. Consider the last period T , any θ^T , and pairs (θ_T, θ_T^*) with $\pi_T(\theta^T, \theta_T^*) = \theta_T$. Since $V_i(T) = \{0\}$, the nestedness of the dynamic lower contour sets, i.e., $L_{i,\theta^T}^f(f(\theta_T), \theta_T) \subseteq L_{i,\theta^T}^{\pi}(f(\theta_T), \theta_T^*)$, is equivalent to the nestedness of the static lower contour sets, i.e., $L_i(f(\theta_T), \theta_T) \subseteq L_i(f(\theta_T), \theta_T^*)$. From Maskin monotonicity, it follows that $f(\theta_T^*) = f(\theta_T)$, as required. To complete the induction argument, consider period $t < T$ and suppose that for all $\tau > t$, for all θ^τ , for all $(\theta_\tau, \theta_\tau^*) \in \Theta \times \Theta$, and for all deceptions π such that $\pi_\tau(\theta^\tau, \theta_\tau^*) = \theta_\tau$, we

	$\theta_L \theta_L \theta_H$	$\theta_L \theta_L \theta_L$	$\theta_L \theta_H \theta_L$	$\theta_L \theta_H \theta_H$	$N \theta_L$	$\theta_H \theta_L \theta_H$	$\theta_H \theta_L \theta_L$	$\theta_H \theta_H \theta_L$	$\theta_H \theta_H \theta_H$	$N \theta_H$
$\theta_L \theta_L \theta_H$	0	0	0	0	$\frac{\delta u(-1) + \delta u(3)}{2}$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$
	$10 + 12\delta$	$10 + 5\delta$	10	$10 + 7\delta$	$10 + 11\delta$	1	1	1	1	1
$\theta_L \theta_L \theta_L$	0	$\frac{\delta u(4)}{2}$	$\frac{\delta u(4)}{2}$	0	$\frac{\delta u(-1) + \delta u(3)}{2}$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$
	$10 + 5\delta$	$10 + 10\delta$	$10 + 5\delta$	10	$10 + 11\delta$	1	1	1	1	1
$\theta_L \theta_H \theta_L$	0	$\frac{\delta u(4)}{2}$	$\frac{\delta u(4) + \delta u(-4)}{2}$	$\frac{\delta u(-4)}{2}$	$\frac{\delta u(-1) + \delta u(3)}{2}$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$
	10	$10 + 5\delta$	$10 + 12\delta$	$10 + 7\delta$	$10 + 11\delta$	1	1	1	1	1
$\theta_L \theta_H \theta_H$	0	0	$\frac{\delta u(-4)}{2}$	$\frac{\delta u(-4)}{2}$	$\frac{\delta u(-1) + \delta u(3)}{2}$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$
	$10 + 7\delta$	10	$10 + 7\delta$	$10 + 14\delta$	$10 + 11\delta$	1	1	1	1	1
$N \theta_L$	$u(-4) + \delta u(y)$	$u(-4) + \delta u(y)$	$u(-4) + \delta u(y)$	$u(-4) + \delta u(y)$	$\frac{(2+\delta)u(-4)}{2}$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$\frac{(2+\delta)u(10) + \delta u(14)}{2}$
	10	10	10	10	$14 + 14\delta$	1	1	1	1	0
$\theta_H \theta_L \theta_H$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-4)$	$u(-4)$	$u(-4)$	$u(-4)$	$\frac{2u(-4) + \delta u(-1) + \delta u(3)}{2}$
	1	1	1	1	1	$14 + 12\delta$	$14 + 5\delta$	14	$14 + 7\delta$	$14 + 11\delta$
$\theta_H \theta_L \theta_L$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-4)$	$\frac{2u(-4) + \delta u(4)}{2}$	$\frac{2u(-4) + \delta u(4)}{2}$	$u(-4)$	$\frac{2u(-4) + \delta u(-1) + \delta u(3)}{2}$
	1	1	1	1	1	$14 + 5\delta$	$14 + 10\delta$	$14 + 5\delta$	14	$14 + 11\delta$
$\theta_H \theta_H \theta_L$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-4)$	$\frac{2u(-4) + \delta u(4)}{2}$	$\frac{(2+\delta)u(-4) + \delta u(4)}{2}$	$\frac{(2+\delta)u(-4)}{2}$	$\frac{2u(-4) + \delta u(-1) + \delta u(3)}{2}$
	1	1	1	1	1	14	$14 + 5\delta$	$14 + 12\delta$	$14 + 7\delta$	$14 + 11\delta$
$\theta_H \theta_H \theta_H$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-4)$	$u(-4)$	$\frac{(2+\delta)u(-4)}{2}$	$\frac{(2+\delta)u(-4)}{2}$	$\frac{2u(-4) + \delta u(-1) + \delta u(3)}{2}$
	1	1	1	1	1	$14 + 7\delta$	14	$14 + 7\delta$	$14 + 14\delta$	$14 + 11\delta$
$N \theta_H$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$\frac{(2+\delta)u(10) + \delta u(14)}{2}$	$u(-4)$	$u(-4)$	$u(-4)$	$u(-4)$	$\frac{(2+\delta)u(-4)}{2}$
	1	1	1	1	0	10	10	10	10	$14 + 14\delta$

TABLE 6. The reduced strategic-form game: first-period quality θ_L .

	$\theta_L \theta_L \theta_H$	$\theta_L \theta_L \theta_L$	$\theta_L \theta_H \theta_L$	$\theta_L \theta_H \theta_H$	$N \theta_L$	$\theta_H \theta_L \theta_H$	$\theta_H \theta_L \theta_L$	$\theta_H \theta_H \theta_L$	$\theta_H \theta_H \theta_H$	$N \theta_H$
$\theta_L \theta_L \theta_H$	$u(4)$	$u(4)$	$u(4)$	$u(4)$	$\frac{2u(4)+\delta u(-1)+\delta u(3)}{2}$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$
	$10+12\delta$	$10+5\delta$	10	$10+7\delta$	$10+11\delta$	1	1	1	1	1
$\theta_L \theta_L \theta_L$	$u(4)$	$\frac{(2+\delta)u(4)}{2}$	$\frac{(2+\delta)u(4)}{2}$	$u(4)$	$\frac{2u(4)+\delta u(-1)+\delta u(3)}{2}$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$
	$10+5\delta$	$10+10\delta$	$10+5\delta$	10	$10+11\delta$	1	1	1	1	1
$\theta_L \theta_H \theta_L$	$u(4)$	$\frac{(2+\delta)u(4)}{2}$	$\frac{(2+\delta)u(4)+\delta u(-4)}{2}$	$\frac{2u(4)+\delta u(-4)}{2}$	$\frac{2u(4)+\delta u(-1)+\delta u(3)}{2}$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$
	10	$10+5\delta$	$10+12\delta$	$10+7\delta$	$10+11\delta$	1	1	1	1	1
$\theta_L \theta_H \theta_H$	$u(4)$	$u(4)$	$\frac{2u(4)+\delta u(-4)}{2}$	$\frac{2u(4)+\delta u(-4)}{2}$	$\frac{2u(4)+\delta u(-1)+\delta u(3)}{2}$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$
	$10+7\delta$	10	$10+7\delta$	$10+14\delta$	$10+11\delta$	1	1	1	1	1
$N \theta_L$	$\delta u(y)$	$\delta u(y)$	$\delta u(y)$	$\delta u(y)$	$\frac{\delta u(-4)}{2}$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$\frac{(2+\delta)u(14)+\delta u(10)}{2}$
	10	10	10	10	$14+14\delta$	1	1	1	1	0
$\theta_H \theta_L \theta_H$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	0	0	0	0	$\frac{\delta u(-1)+\delta u(3)}{2}$
	1	1	1	1	1	$14+12\delta$	$14+5\delta$	14	$14+7\delta$	$14+11\delta$
$\theta_H \theta_L \theta_L$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	0	$\frac{\delta u(4)}{2}$	$\frac{\delta u(4)}{2}$	0	$\frac{\delta u(-1)+\delta u(3)}{2}$
	1	1	1	1	1	$14+5\delta$	$14+10\delta$	$14+5\delta$	14	$14+11\delta$
$\theta_H \theta_H \theta_L$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	0	$\frac{\delta u(4)}{2}$	$\frac{\delta u(-4)+\delta u(4)}{2}$	$\frac{\delta u(-4)}{2}$	$\frac{\delta u(-1)+\delta u(3)}{2}$
	1	1	1	1	1	14	$14+5\delta$	$14+12\delta$	$14+7\delta$	$14+11\delta$
$\theta_H \theta_H \theta_H$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	0	0	$\frac{\delta u(-4)}{2}$	$\frac{\delta u(-4)}{2}$	$\frac{\delta u(-1)+\delta u(3)}{2}$
	1	1	1	1	1	$14+7\delta$	14	$14+7\delta$	$14+14\delta$	$14+11\delta$
$N \theta_H$	$u(-1)$	$u(-1)$	$u(-1)$	$u(-1)$	$\frac{(2+\delta)u(14)+\delta u(10)}{2}$	0	0	0	0	$\frac{\delta u(-4)}{2}$
	1	1	1	1	0	10	10	10	10	$14+14\delta$

TABLE 7. The reduced strategic-form game: first-period quality θ_H .

have that $f(\theta_t^*) = f(\theta_t)$. It follows that in period t the continuation payoff $v_i^{f_\pi}(\theta^t, \theta_t)$ is equal to $v_i^f(t)$ for all agents i , for all (θ^t, θ_t) and, thus, $L_{i,\theta^t}^{f_\pi}(f(\theta_t), \theta_t^*) = L_{i,t}^f(f(\theta_t), \theta_t^*)$ for all $(\theta^t, \theta_t, \theta_t^*)$. As a result, $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,\theta^t}^{f_\pi}(f(\theta_t), \theta_t^*)$ is equivalent to $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,t}^f(f(\theta_t), \theta_t^*)$. In turn, this is equivalent to the nestedness of the static lower contour sets, i.e., $L_i(f(\theta_t), \theta_t) \subseteq L_i(f(\theta_t), \theta_t^*)$. Maskin monotonicity then implies $f(\theta_t^*) = f(\theta_t)$. This concludes the proof. $\square$

PROOF OF REMARK 3. Assume to the contrary that f is dynamic monotonic but not weakly efficient in the range; that is, there exists $\varepsilon > 0$ and a payoff profile $(v_i)_{i \in \mathcal{I}} \in \text{co}(V(f))$ such that $v_i > v_i^f + 2\varepsilon$ for all $i \in \mathcal{I}$. Using a standard argument about convexifying the set of payoffs without public randomization (e.g., see Lemma 3.7.2 in Mailath and Samuelson 2006), it follows that there exists δ_{H_2} such that for all $\delta \in (\delta_{H_2}, 1)$ there exists an infinite sequence of social choice functions $\{f_1, f_2, \dots\}$ with $f_t \in \mathcal{F}(f)$ for all integers t (i.e., the range of f_t is a subset of the range of f), and $(1 - \delta) \sum_{\tau=t}^{\infty} \delta^{\tau-t} v_i^{f_\tau} > v_i - \varepsilon$, for all $i \in \mathcal{I}$, for all t . Since $f_t \in \mathcal{F}(f)$, there exist mappings $\pi'_t : \Theta \rightarrow \Theta$ such that $f \circ \pi'_t = f_t$. Consider the deception π such that $\pi_t(\theta^t, \theta_t^*) = \pi'_t(\theta_t^*)$ for all θ_t^* , for all θ^t , for all t . It follows that $v_i^{f_\pi}(\theta^t, \theta_t^*) = (1 - \delta) \sum_{\tau=t+1}^{\infty} \delta^{\tau-t-1} v_i^{f_\tau} > v_i - \varepsilon > v_i^f + \varepsilon$ for all i . Let $\rho = \max_{i \in \mathcal{I}, \theta, \theta^* \in \Theta} |u_i(f(\theta), \theta) - u_i(f(\theta), \theta^*)|$, and let $\delta_H = \max(\rho/(\rho + \varepsilon), \delta_{H_2})$. Then, for $\delta \in (\delta_H, 1)$, for all i , for all pairs (θ_t, θ_t^*) with $\pi_t(\theta^t, \theta_t^*) = \theta_t$, for all $\theta^t \in \Theta^t$, for all $t \in \mathcal{T}$, it is $(1 - \delta)u_i(f(\theta_t), \theta_t^*) + \delta v_i^{f_\pi}(\theta^t, \theta_t^*) \geq (1 - \delta)u_i(f(\theta_t), \theta_t) + \delta v_i^f$ or, equivalently, $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,\theta^t}^{f_\pi}(f(\theta_t), \theta_t^*)$. Dynamic monotonicity then implies that $f \circ \pi'_t = f_t = f$ for all t , contradicting the assumed weak inefficiency of f . $\square$

PROOF OF REMARK 4. Assume f is Maskin monotonic and efficient in the range, and suppose that there exists a deception π such that for all $t \in \mathcal{T}$, for all $\theta^t \in \Theta^t$, for all pairs (θ_t, θ_t^*) with $\pi_t(\theta^t, \theta_t^*) = \theta_t$, we have $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,\theta^t}^{f_\pi}(f(\theta_t), \theta_t^*)$ for all $i \in \mathcal{I}$. Recall that $v_i^{f_\pi}(\theta^t, \theta_t)$ is the (normalized) expected discounted continuation payoff of agent i from following the deception π from the history induced by π and (θ^t, θ_t) . Thus, $v_i^{f_\pi}(\theta^t, \theta_t)$ is an element of the convex hull of $V(f)$, the set of payoff profiles of social choice functions with a range contained in the range of f . First, suppose that $(v_i^{f_\pi}(\theta^t, \theta_t))_{i \in \mathcal{I}} \neq (v_i^f)_{i \in \mathcal{I}}$. Since f is efficient in the range, it follows that there exists an agent i^* such that $v_{i^*}^{f_\pi}(\theta^t, \theta_t) < v_{i^*}^f$. Consequently, we have that $(f(\theta_t), v_{i^*}^f) \in L_{i^*,t}^f(f(\theta_t), \theta_t)$ (by definition) and $(f(\theta_t), v_{i^*}^f) \notin L_{i^*,\theta^t}^{f_\pi}(f(\theta_t), \theta_t^*)$, a contradiction. So it must be that $v_i^{f_\pi}(\theta^t, \theta_t) = v_i^f$ for all $i \in \mathcal{I}$. It then immediately follows that the nestedness of the dynamic lower contour sets (i.e., $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,\theta^t}^{f_\pi}(f(\theta_t), \theta_t^*)$) implies the nestedness of the static lower contour sets (i.e., $L_i(f(\theta_t), \theta_t) \subseteq L_i(f(\theta_t), \theta_t^*)$). Maskin monotonicity then implies that $f(\theta_t^*) = f(\theta_t)$. This shows that $f(\pi_t(\theta^t, \cdot)) = f$ for all $\theta^t \in \Theta^t$, for all $t \in \mathcal{T}$, and hence f must be dynamic monotonic. $\square$

PROOF OF REMARK 6. By contradiction, suppose that f is dynamic monotonic over T periods, but not over $T + 1$ periods. Since f is not dynamic monotonic over $T + 1$ periods, there exist a profile of states $\theta^{T+1} \in \Theta^{T+1}$ and a deception $\pi \in \Pi^{T+1}$ with $f(\pi_t(\theta^t, \theta_t)) \neq$

$f(\theta_t)$ for at least one $t \in \{1, \dots, T+1\}$, while the dynamic lower contour sets are nested, i.e., for all $i \in \mathcal{I}$, for all $t \in \mathcal{T}$,

$$L_{i,t}^f(f(\pi_t(\theta^t, \theta_t)), \pi_t(\theta^t, \theta_t)) \subseteq L_{i,\theta^t}^{f_\pi}(f(\pi_t(\theta^t, \theta_t)), \theta_t). \quad (2)$$

We first argue that $f(\pi_t(\theta^t, \theta_t)) = f(\theta_t)$ for all $t \in \{2, \dots, T+1\}$. Fix the first-period state θ_1 in the profile θ^{T+1} and consider any deception $\pi^{**} \in \Pi^T$ such that $\pi_t^{**}(\theta^t, \theta_t) = \pi_{t+1}((\theta_1, \theta^t), \theta_t)$ for all $t \in \{1, \dots, T\}$. In words, π^{**} mirrors the last T periods of π , given that the first-period state was θ_1 .

By (2) and $\beta_{t,T} = \beta_{t+1,T+1}$, for all $i \in \mathcal{I}$ and $t \in \{1, \dots, T\}$,

$$L_{i,t}^f(f(\pi_t^{**}(\theta^t, \theta_t)), \pi_t^{**}(\theta^t, \theta_t)) \subseteq L_{i,\theta^t}^{f_{\pi^{**}}}(f(\pi_t^{**}(\theta^t, \theta_t)), \theta_t).$$

Since f is dynamic monotonic over T periods, this implies that $f(\pi_t^{**}(\theta^t, \theta_t)) = f(\theta_t)$ for all $t \in \{1, \dots, T\}$ or, equivalently, $f(\pi_t(\theta^t, \theta_t)) = f(\theta_t)$ for all $t \in \{2, \dots, T+1\}$. It follows that $v_i^{f_\pi}((\theta^t, \theta_t)) = v_i^{f_{\pi^{**}}}((\theta^t, \theta_t)) = v_i^f(t)$ for all $t \geq 1$.

Therefore, we must have $f(\pi_1(\theta_1)) \neq f(\theta_1)$. We now argue that this cannot be the case either. Consider any deception π° such that $\pi_t^\circ(\theta^t, \theta_t) = \pi_t(\theta^t, \theta_t)$ for all (θ^t, θ_t) , that is, π° coincides with the first T periods of π .

Since f is dynamic monotonic over T periods (and the fact that $f(\pi_1^\circ(\theta_1)) \neq f(\theta_1)$), there exist $i \in \mathcal{I}$, $t \in \{1, \dots, T\}$, and (x, v_i) such that

$$(1 - \beta_{t,T})u_i(f(\pi_t^\circ(\theta^t, \theta_t)), \pi_t^\circ(\theta^t, \theta_t)) + \beta_{t,T}v_i^f(t) \geq (1 - \beta_{t,T})u_i(x, \pi_t^\circ(\theta^t, \theta_t)) + \beta_{t,T}v_i$$

and

$$(1 - \beta_{t,T})u_i(f(\pi_t^\circ(\theta^t, \theta_t)), \theta_t) + \beta_{t,T}v_i^f(t) < (1 - \beta_{t,T})u_i(x, \theta_t) + \beta_{t,T}v_i.$$

Using the definition of π° , this is equivalent to (remember that $\beta_{t,T+1} \in (0, 1)$)

$$\begin{aligned} (1 - \beta_{t,T+1})[u_i(f(\pi_t(\theta^t, \theta_t)), \pi_t(\theta^t, \theta_t)) - u_i(x, \pi_t(\theta^t, \theta_t))] \\ \geq \frac{\beta_{t,T}}{1 - \beta_{t,T}}(1 - \beta_{t,T+1})[v_i - v_i^f(t)] \\ > (1 - \beta_{t,T+1})[u_i(f(\pi_t(\theta^t, \theta_t)), \theta_t) - u_i(x, \theta_t)]. \end{aligned}$$

Let $\hat{v}_i$ be given by

$$\frac{\beta_{t,T}}{1 - \beta_{t,T}} \frac{1 - \beta_{t,T+1}}{\beta_{t,T+1}} v_i + \left(1 - \frac{\beta_{t,T}}{1 - \beta_{t,T}} \frac{1 - \beta_{t,T+1}}{\beta_{t,T+1}}\right) v_i^f(t).$$

Since $\beta_{t,T} \leq \beta_{t,T+1}$, we have that $\hat{v}_i \in [\underline{v}_i, \bar{v}_i]$. It follows that there exists $(x, \hat{v}_i) \in X \times V_i(t)$ such that

$$\begin{aligned} (1 - \beta_{t,T+1})u_i(f(\pi_t(\theta^t, \theta_t)), \pi_t(\theta^t, \theta_t)) + \beta_{t,T+1}v_i^f(t) \\ \geq (1 - \beta_{t,T+1})u_i(x, \pi_t(\theta^t, \theta_t)) + \beta_{t,T+1}\hat{v}_i \end{aligned}$$

and

$$(1 - \beta_{t,T+1})u_i(f(\pi_t(\theta^t, \theta_t)), \theta_t) + \beta_{t,T+1}v_i^f(t) < (1 - \beta_{t,T+1})u_i(x, \theta_t) + \beta_{t,T+1}\hat{v}_i.$$

This is equivalent to $L_{i,t}^f(f(\pi_t(\theta^t, \theta_t)), \pi_t(\theta^t, \theta_t)) \not\subseteq L_{i,\theta^t}^{f_\pi}(f(\pi_t(\theta^t, \theta_t)), \theta_t)$, a contradiction with (2). Therefore, $f(\pi_1(\theta_1)) = f(\theta_1)$, as required. $\square$

PROOF OF THEOREM 1. Suppose that f is repeatedly implementable by the dynamic mechanism regime r . Fix an equilibrium s . Consider a history h^t and a mechanism $G_t = \langle M^{G_t}, g_t \rangle$ having positive probability of occurring on the equilibrium path at period t ; that is, such that $q(h^t; s) > 0$ and $r(G_t; h^t) > 0$. Since the dynamic regime r implements f , the profile of actions $s(h^t, G_t, \theta_t)$ at period t must satisfy $g_t(s(h^t, G_t, \theta_t)) = f(\theta_t)$ for each $\theta_t \in \Theta$, and the continuation payoff must be $v_i^f(t)$. Let $Q_i(h^t, G_t, \theta_t; s)$ be the set of current alternative and continuation payoff pairs that agent i is able to generate by any deviation starting at t , given that all other agents follow s . Formally, $(x, v_i) \in X \times V_i(t)$ belongs to $Q_i(h^t, G_t, \theta_t; s)$ if there exists $m_i \in M_i^{G_t}$ such that $x = g(m_i, s_{-i}(h^t, G_t, \theta_t))$ and there exists $v_i \in V_i(t)$ that corresponds to i 's expected discounted continuation payoff when (starting at t , in state θ_t , after history h^t) agent i follows some continuation strategy (which prescribes sending message m_i at t), while all other agents continue to follow s_{-i} .

Since s is an equilibrium, for each $i \in \mathcal{I}$, for each $\theta_t \in \Theta$, we must have that

$$(1 - \beta_{t,T})u_i(f(\theta_t), \theta_t) + \beta_{t,T}v_i^f(t) \geq (1 - \beta_{t,T})u_i(x, \theta_t) + \beta_{t,T}v_i$$

for each $(x, v_i) \in Q_i(h^t, G_t, \theta_t; s)$. Consequently, it must be that $Q_i(h^t, G_t, \theta_t; s) \subseteq L_{i,t}^f(f(\theta_t), \theta_t)$ for all h^t, θ_t , and G_t such that $q(h^t; s) > 0$ and $r(G_t; h^t) > 0$.

Now consider a deception π such that for all $t \in \mathcal{T}$, for all $\theta^t \in \Theta^t$, for all pairs (θ_t, θ_t^*) with $\pi_t(\theta^t, \theta_t^*) = \theta_t$, we have $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,\theta^t}^{f_\pi}(f(\theta_t), \theta_t^*)$ for all $i \in \mathcal{I}$. In the remainder of the proof, we will show that there exists an equilibrium s' that implements the social choice function $f(\pi_t(\theta^t, \cdot))$ at each period t for each θ^t . Since the regime r repeatedly implements f , it must be that $f(\pi_t(\theta^t, \cdot)) = f$ for all $\theta^t \in \Theta^t$, for all $t \in \mathcal{T}$. Hence, we may conclude that f is dynamic monotonic and the theorem holds.

We now construct the strategy profile s' . First, consider the equilibrium path. Let $h^1 = h_\pi^1 = \{\emptyset\}$ and for all θ_1 , for all G_1 , for all i , define

$$s'_i(h^1, G_1, \theta_1) = s_i(h_\pi^1, G_1, \pi_1(\theta_1)).$$

Then assume that the strategy profile s' and the histories h^τ and h_π^τ have been defined up to period $\tau = t$. Let $h^{t+1} = (h^t, G_t, \theta_t, s'(h^t, G_t, \theta_t))$, with $h^t = (h_D^t, \theta^t)$, be a period $t + 1$ history corresponding to the history of realized states θ^t . Associate the history h^{t+1} with $h_\pi^{t+1} = (h_\pi^t, G_t, \pi_t(\theta^t, \theta_t), s(h_\pi^t, G_t, \pi_t(\theta^t, \theta_t)))$, and for all θ_{t+1} , for all G_{t+1} , for all i , define

$$s'_i(h^{t+1}, G_{t+1}, \theta_{t+1}) = s_i(h_\pi^{t+1}, G_{t+1}, \pi_{t+1}((\theta^t, \theta_t), \theta_{t+1})).$$

This concludes the definition of s' on the equilibrium path. Note that it prescribes that agents behave as s would prescribe if the history of realized states were the one described by the deception π instead of the true history.

We now define s' when agent i unilaterally deviates from the equilibrium path at period t , history $h^t = (h_D^t, \theta^t)$, and state θ_t . The history induced by deviating to $m_{i,t} \neq s'_i(h^t, G_t, \theta_t)$ is $h^{t+1}|_{m_{i,t}} = (h^t, G_t, \theta_t, (m_{i,t}, s'_{-i}(h^t, G_t, \theta_t)))$. Associate $h^{t+1}|_{m_{i,t}}$ with $h^{t+1}|_{m_{i,t}} = (h_{\pi}^t, G_t, \pi_t(\theta^t, \theta_t), (m_{i,t}, s_{-i}(h_{\pi}^t, G_t, \pi_t(\theta^t, \theta_t))))$. For all θ_{t+1} , for all G_{t+1} , for all i , define

$$s'_i(h^{t+1}|_{m_{i,t}}, G_{t+1}, \theta_{t+1}) = s_i(h_{\pi}^{t+1}|_{m_{i,t}}, G_{t+1}, \theta_{t+1}).$$

Decompose history $h^{t+\tau}$ into the history up to $t+1$, h^{t+1} , and the history after $t+1$, $h^{t+1,t+\tau}$, and write $h^{t+\tau} = (h^{t+1}, h^{t+1,t+\tau})$. For all $\tau \geq 2$, for all histories $h^{t+\tau} = (h^{t+1}|_{m_{i,t}}, h^{t+1,t+\tau})$, define

$$s'_i((h^{t+1}|_{m_{i,t}}, h^{t+1,t+\tau}), G_{t+\tau}, \theta_{t+\tau}) = s_i((h_{\pi}^{t+1}|_{m_{i,t}}, h^{t+1,t+\tau}), G_{t+\tau}, \theta_{t+\tau})$$

for all $\theta_{t+\tau}$, for all $G_{t+\tau}$, for all i . Finally, assume that s' agrees with s at all other histories. Note that s' prescribes that following the deviation by agent i , starting from period $t+1$, agents revert to the original equilibrium strategy profile s .

By construction of s' , for any t , any θ^t , and any θ_t^* , the expected payoff of agent i at state θ_t^* from period t onward is

$$(1 - \beta_{t,T})u_i(f(\pi_t(\theta^t, \theta_t^*)), \theta_t^*) + \beta_{t,T}v_i^{f_{\pi}}(\theta^t, \theta_t^*).$$

In addition, if agent i deviates from s' at history $((h_D^t, \theta^t), G_t, \theta_t^*)$ by announcing $m_{i,t}$, the alternative implemented is x satisfying

$$x = g(m_{i,t}, s'_{-i}((h_D^t, \theta^t), G_t, \theta_t^*)) = g(m_{i,t}, s_{-i}(h_{\pi}^t, G_t, \pi_t(\theta^t, \theta_t^*)),$$

and i 's continuation payoff v_i must satisfy $(x, v_i) \in Q_i(h_{\pi}^t, G_t, \theta_t; s)$, where $\theta_t = \pi_t(\theta^t, \theta_t^*)$. Thus, if agent i has a profitable deviation, there exist an alternative x and a continuation payoff v_i such that $(x, v_i) \in Q_i(h_{\pi}^t, G_t, \theta_t; s)$ and

$$(1 - \beta_{t,T})u_i(x, \theta_t^*) + \beta_{t,T}v_i > (1 - \beta_{t,T})u_i(f(\pi_t(\theta^t, \theta_t^*)), \theta_t^*) + \beta_{t,T}v_i^{f_{\pi}}(\theta^t, \theta_t^*),$$

or, since $f(\pi_t(\theta^t, \theta_t^*)) = f(\theta_t)$, $(x, v_i) \notin L_{i,\theta_t}^{f_{\pi}}(f(\theta_t), \theta_t^*)$.

By construction, the t -period deviation by i is feasible under strategy profile s when the state is $\theta_t = \pi_t(\theta^t, \theta_t^*)$ and the history is h_{π}^t . Since, by assumption, $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,\theta_t}^{f_{\pi}}(f(\theta_t), \theta_t^*)$, it must be that $(x, v_i) \notin L_{i,t}^f(f(\theta_t), \theta_t)$ and hence the deviation from s at t is profitable. This contradicts the assumption that s is an equilibrium. Hence, it cannot be $(x, v_i) \in Q_i(h_{\pi}^t, G_t, \theta_t; s)$ and $(x, v_i) \notin L_{i,\theta_t}^{f_{\pi}}(f(\theta_t), \theta_t^*)$; it must be $Q_i(h_{\pi}^t, G_t, \theta_t; s) \subseteq L_{i,\theta_t}^{f_{\pi}}(f(\theta_t), \theta_t^*)$. It follows that s' is an equilibrium (no agent has a profitable deviation at any point in time). Since the mechanism regime r repeatedly implements f , it must therefore be that $f(\pi_t(\theta^t, \cdot)) = f$ for all $\theta^t \in \Theta^t$, for all $t \in \mathcal{T}$. This concludes the proof of the necessity of dynamic monotonicity. $\square$

PROOF OF THEOREM 2. Assume that the social choice function f is dynamic monotonic and satisfies no-veto power. We show that f is repeatedly implementable.

Step 1: Static mechanisms. We present several static mechanisms.

◇ *The period t canonical mechanism, $G_t^* = \langle M_t^*, g_t^* \rangle$.* Let $\mathbb{N}$ be the set of nonnegative integers. For each $i \in \mathcal{I}$, the message space of agent i is $M_{t,i}^* = \Theta \times X \times V_i(t) \times \mathbb{N}$, with $m_{t,i} = (\theta_{t,i}, x_{t,i}, v_{t,i}, n_{t,i})$ a generic element. The allocation rule g_t^* is defined as follows:

Rule 1.1. If $m_{t,i} = (\theta_t, f(\theta_t), v_i^f(t), 0)$ for each $i \in \mathcal{I}$, then $g_t^*(m_{t,1}, \dots, m_{t,I}) = f(\theta_t)$.

Rule 1.2. If there exists j such that $m_{t,i} = (\theta_t, f(\theta_t), v_i^f(t), 0)$ for each $i \in \mathcal{I} \setminus \{j\}$ and $m_{t,j} = (\theta_{t,j}, x_{t,j}, v_{t,j}, n_{t,j}) \neq (\theta_t, f(\theta_t), v_j^f(t), 0)$, then $g_t^*(m_{t,1}, \dots, m_{t,I}) = x_{t,j}$ if $(x_{t,j}, v_{t,j}) \in L_{j,t}^f(f(\theta_t), \theta_t)$, and $g_t^*(m_{t,1}, \dots, m_{t,I}) = f(\theta_t)$ otherwise.

Rule 1.3. If neither Rule 1.1 nor Rule 1.2 applies, then $g_t^*(m_{t,1}, \dots, m_{t,I}) = x_{t,i^*}$ with $i^* = \min\{i \in \mathcal{I} : n_{t,i} \geq n_{t,j} \text{ for all } j \in \mathcal{I}\}$.

◇ *Agent i 's dictatorship, $D_i = \langle M^{D_i}, g^{D_i} \rangle$.* The agents' message spaces are $M_i^{D_i} = X$ and $M_j^{D_i} = \{\emptyset\}$ for $j \in \mathcal{I} \setminus \{i\}$. The allocation rule is $g^{D_i}(m_i, m_{-i}) = m_i$.

◇ *The "punishment" mechanism, $P_i = \langle M^{P_i}, g^{P_i} \rangle$.* The message space is $M_j^{P_i} = X$ for all $j \in \mathcal{I}$. If for all $j \in \mathcal{I}^*$ with $|\mathcal{I}^*| \geq n - 1$, $m_j = x$, then the allocation rule is $g^{P_i}((m_j)_{j \in \mathcal{I}}) = x$; otherwise, $g^{P_i}((m_j)_{j \in \mathcal{I}}) = m_{i+1 \pmod{I}}$.

Step 2: The dynamic mechanism regime r . We define the transition probability $r(G_t, h_D^t)$ that, after the designer history h_D^t , the mechanism in period t is G_t .

Period 1. At the initial history, the mechanism is G_1^* ; that is, $r(G_1^*; \emptyset) = 1$.

Period t .

(A) Suppose that the history at period t is $h_D^t = (h_D^{t-1}, G_{t-1}^*, (m_{t-1,i})_{i \in \mathcal{I}})$ (i.e., the mechanism was G_{t-1}^* in period $t - 1$). The transition to period t is as follows:

- If $m_{t-1,i} = (\theta_{t-1}, f(\theta_{t-1}), v_i^f(t-1), 0)$ for each $i \in \mathcal{I}$ and some $\theta_{t-1} \in \Theta$, then $r(G_t^*; h_D^t) = 1$. In words, if Rule 1.1 of G_{t-1}^* applied in period $t - 1$, then in period t the mechanism is G_t^* with probability 1.
- If there exists j such that $m_{t-1,i} = (\theta_{t-1}, f(\theta_{t-1}), v_i^f(t-1), 0)$ for each $i \in \mathcal{I} \setminus \{j\}$ and $m_{t-1,j} = (\theta_{t-1,j}, x_{t-1,j}, v_{t-1,j}, n_{t-1,j}) \neq (\theta_{t-1}, f(\theta_{t-1}), v_j^f(t-1), 0)$, then $r(P_j; h_D^t) = (1 - \lambda_j^{(t)})$ and $r(D_j; h_D^t) = \lambda_j^{(t)}$. In words, if Rule 1.2 of G_{t-1}^* applied in period $t - 1$ with j as the odd man out, then the mechanism in period t is the "punishment" mechanism P_j with probability $(1 - \lambda_j^{(t)})$ and the dictatorial mechanism D_j with probability $\lambda_j^{(t)}$ (to be defined later). As we shall see in (B) and (C), once either P_j or D_j is selected at t , it is adopted in all future periods.
- If any other profile of messages is played in period $t - 1$, then $r(D_{i^*}; h_D^t) = 1$; that is, the period t mechanism is D_{i^*} with i^* the lowest indexed agent having announced the highest integer in period $t - 1$.

(B) If the designer history at period t is $h_D^t = (h_D^{t-1}, D_j, (m_{t-1,i})_{i \in \mathcal{I}})$ (i.e., the mechanism at period $t - 1$ was D_j), then $r(D_j; h_D^t) = 1$.

(C) If the designer history at period t is $h_D^t = (h_D^{t-1}, P_j, (m_{t-1,i})_{i \in \mathcal{I}})$ (i.e., the mechanism at period $t - 1$ was P_j), then $r(P_j; h_D^t) = 1$.

Step 3: Definition of $\lambda_j^{(t)}$, $t \in T \setminus \{1\}$. We only need to define $\lambda_j^{(t)}$ when Rule 1.2 of G_{t-1}^* applied at $t - 1$ and j was the odd man out. Let $(x_{t-1,j}, v_{t-1,j})$ be the pair of alternative and continuation payoff announced by j in period $t - 1$. Recall that $\underline{v}_j$ and $\bar{v}_j$ are the lowest and highest expected payoffs agent j can obtain. If $\bar{v}_j = \underline{v}_j$, we may choose any $\lambda_j^{(t)} \in [0, 1]$. If $\bar{v}_j > \underline{v}_j$, define $\lambda_j^{(t)} \in [0, 1]$ as the unique solution of

$$v_{t-1,j} = \lambda_j^{(t)} \bar{v}_j + (1 - \lambda_j^{(t)}) \underline{v}_j$$

if $(x_{t-1,j}, v_{t-1,j}) \in L_{j,t}^f(f(\theta_{t-1}), \theta_{t-1})$, and otherwise as the unique solution of

$$v_j^f = \lambda_j^{(t)} \bar{v}_j + (1 - \lambda_j^{(t)}) \underline{v}_j.$$

Step 4: Existence of an equilibrium. There exists an equilibrium s^E that repeatedly implements f .

For each agent i , the strategy s_i^E is defined as follows:

- 4.1. For all $\theta_1 \in \Theta$, $s_i^E(\emptyset, G_1^*, \theta_1) = (\theta_1, f(\theta_1), v_i^f(1), 0)$.
- 4.2. For all $\theta_t \in \Theta$, if h^t is such that for all $\tau < t$, (i) $G_\tau = G_\tau^*$ and (ii) $m_\tau = (\theta_\tau, f(\theta_\tau), v_i^f(\tau), 0)_{i \in \mathcal{I}}$ for some $\theta_\tau \in \Theta$, then $s_i^E(h^t, G_t^*, \theta_t) = (\theta_t, f(\theta_t), v_i^f(t), 0)$.
- 4.3. For all $h^t \in \mathcal{H}$, all $\theta_t \in \Theta$, and all $i \in \mathcal{I}$, $s_i^E(h^t, P_j, \theta_t) = \underline{x}_j^{\theta_t}$, where $\underline{x}_j^{\theta_t} \in \arg \min_{x \in X} u_j(x, \theta_t)$.
- 4.4. For all $h^t \in \mathcal{H}$ and all $\theta_t \in \Theta$, $s_i^E(h^t, D_j, \theta_t) = \emptyset$ if $i \neq j$, and $s_j^E(h^t, D_j, \theta_t) = \bar{x}_j^{\theta_t}$ with $\bar{x}_j^{\theta_t} \in \arg \max_{x \in X} u_j(x, \theta_t)$.

According to s_i^E , in the first period, each agent i announces $(\theta_1, f(\theta_1), v_i^f(1), 0)$ whenever θ_1 is the true state. In period $t > 1$, there are three cases. First, if the game being played is G_t^* and all agents have made “unanimous” announcements $(\theta_\tau, f(\theta_\tau), v_i^f(\tau), 0)$ in all past periods $\tau < t$, then agent i announces $(\theta_t, f(\theta_t), v_i^f(t), 0)$ whenever θ_t is the true state in period t . Second, if the game being played is P_j , then *all* agents announce an alternative that “min-max” agent j . Third, if the game being played is D_j , then agent j chooses an alternative that maximizes his period t payoff.

Under s^E , agent j 's expected payoff at period t when the state is θ_t is

$$(1 - \beta_{t,T})u_j(f(\theta_t), \theta_t) + \beta_{t,T}v_j^f(t).$$

If agent j deviates and announces $(\theta_{t,j}, x_{t,j}, v_{t,j}, n_{t,j}) \neq (\theta_t, f(\theta_t), v_j^f(t), 0)$, the high-est possible payoff following the deviation is

$$\min\{(1 - \beta_{t,T})u_j(x_{t,j}, \theta_t) + \beta_{t,T}v_{t,j}, (1 - \beta_{t,T})u_j(f(\theta_t), \theta_t) + \beta_{t,T}v_j^f(t)\},$$

so that agent j has no profitable deviation. Note that agent j obtains a continuation payoff of $v_{t,j}$ if, following the deviation, he announces $x_{\tau,j} \in \arg \max_{x \in X} u_i(x, \theta_\tau)$ for all $\tau > t$, for all $\theta_\tau \in \Theta$, whenever he is dictatorial.

Step 5: No undesirable equilibria. There are no undesirable equilibria.

Let s be any equilibrium and consider any history h^t with $q(h^t; s) > 0$. We want to show that (i) $g(s(h^t, G_i^*, \theta_t)) = f(\theta_t)$ for all $\theta_t \in \Theta$ if $r(G_i^*; h^t) > 0$, (ii) $g(s(h^t, D_i, \theta_t)) = f(\theta_t)$ for all $\theta_t \in \Theta$ if $r(D_i; h^t) > 0$, and (iii) $g(s(h^t, P_i, \theta_t)) = f(\theta_t)$ for all $\theta_t \in \Theta$ if $r(P_i; h^t) > 0$.

Statements (ii) and (iii) follow from no-veto power. If $G_i \in \{D_i, P_i\}$ is adopted in period t with positive probability, then there was a last time $t' < t$ when G_i^* was played and either Rule 1.2 with agent i the odd man out or Rule 1.3 with i the dictator, applied. Every agent j other than agent i could have deviated and become the dictator at t' and in all future periods. For such a deviation not to be profitable it must be the case that the alternative implemented at t and state θ_t is $x \in \max_j^{\theta_t} X$ for all $j \neq i$. No-veto power then implies $x = f(\theta_t)$, as statements (ii) and (iii) claim.²²

Now consider statement (i). Assume that $r(G_i^*; h^t) > 0$.

CLAIM 1. *If the equilibrium s is such that $s(h^t, G_i^*, \theta_t)$ corresponds to Rule 1.2 of G_i^* , i.e., $s_i(h^t, G_i^*, \theta_t) = (\tilde{\theta}_t, f(\tilde{\theta}_t), v_i^f(t), 0)$ for each $i \in \mathcal{I} \setminus \{j\}$ and $s_j(h^t, G_i^*, \theta_t) = (\theta_{t,j}, x_{t,j}, v_{t,j}, n_{t,j}) \neq (\tilde{\theta}_t, f(\tilde{\theta}_t), v_i^f(t), 0)$, then the alternative implemented at θ_t is $f(\theta_t)$.*

PROOF. Let x be the alternative implemented. Note that since Rule 1.2 of G_i^* applies, x is either $x_{t,j}$ or $f(\tilde{\theta}_t)$. At the history (h^t, G_i^*, θ_t) , any agent $i \neq j$ can deviate and announce $(\theta_{t,i}, x_i^{\theta_t}, v_{t,i}, n_{t,i})$, with $n_{t,i} > n_{t,j}$, $x_i^{\theta_t} \in \max_i^{\theta_t} X$, and then choose $x_i^{\theta_\tau} \in \max_i^{\theta_\tau} X$ when the mechanism D_i is played in period τ and θ_τ is the realized state, for any $\tau > t$. Since agent i becomes dictator for all $\tau \geq t$, the expected payoff starting at t from such a deviation is $(1 - \beta_{t,T})u_i(x_i^{\theta_t}, \theta_t) + \beta_{t,T}\bar{v}_i$. For the deviation not to be profitable, it must be that $x \in \max_i^{\theta_t} X$ for all $i \in \mathcal{I} \setminus \{j\}$. It follows from no-veto power that $x = f(\theta_t)$.²³ $\triangleleft$

CLAIM 2. *If the equilibrium s is such that $s(h^t, G_i^*, \theta_t)$ corresponds to Rule 1.3 of G_i^* , then the alternative implemented at θ_t is $f(\theta_t)$.*

²² To prove that (ii) and (iii) hold under [Assumption A](#) in place of no-veto power, first consider the case when Rule 1.2 applies in period t' . Since each agent j other than i can deviate at t' and become dictator in all subsequent periods τ including t (and also at t'), for Rule 1.2 at t' to be part of an equilibrium it must be that on the equilibrium path in all periods $\tau > t'$ and in all states θ , the alternative chosen maximizes the payoff of each agent $j \neq i$; that is, it must belong to $\bigcap_{j \neq i} \max_j^{\theta} X$. Write $\varphi_{t'}(\tau, \theta)$ for the alternative implemented at τ in state θ if the mechanism is P_i , and write $\varphi_{t'}(\tau, \theta)$ if the mechanism is D_i . Clearly, it must be that $\varphi_{t'}(\tau, \theta) \in \max_i^{\theta} X$ for all $\tau > t'$, for all θ . Let the second and third elements of the message sent by agent i at t' on the equilibrium path be $(x, v_i(t'))$; at t' agent i must not have a profitable deviation $(y, v_i) \in L_{i,t'}(f(\theta_{t'}), \theta_{t'})$, where $\theta_{t'}$ is the state announced by all agents other than i . The most severe punishment that the other agents could use when mechanism P_i is played at $t > t'$ after a deviation yields agent i a continuation payoff $\underline{v}_i$; the highest payoff that agent i could secure himself after a deviation when mechanism D_i is played at $t > t'$ is $\bar{v}_i$. Since $\lambda(v_i(t'))$ is the probability that D_i is played on the equilibrium path and $\lambda(v_i)$ is the probability D_i is played after the deviation, for the equilibrium under Rule 1.2 at t' to exist it must be that $\beta_{t',T}u_i(x, \theta_{t'}^*) + (1 - \beta_{t',T})[\lambda(v_i(t'))v_i^{\varphi_{t'}} + (1 - \lambda(v_i(t')))\bar{v}_i] \geq \beta_{t',T}u_i(y, \theta_{t'}^*) + (1 - \beta_{t',T})[\lambda(v_i)\bar{v}_i + (1 - \lambda(v_i))\underline{v}_i] = \beta_{t',T}u_i(y, \theta_{t'}^*) + (1 - \beta_{t',T})v_i$ for all $(y, v_i) \in L_{i,t'}(f(\theta_{t'}), \theta_{t'})$, where $\theta_{t'}$ is the state reported by all agents other than i and $\theta_{t'}^*$ is the true state at t' . The result follows from (i) of [Assumption A](#), since $\lambda(v_i(t')) \neq 0$ implies $\varphi_{t'}(\tau, \theta) \in \bigcap_j \max_j^{\theta} X$ for all $\theta \in \Theta$, $\tau > t'$.

Second, if Rule 1.3 applies at t' , the result immediately follows from condition (ii) of [Assumption A](#).

²³The proof that [Claim 1](#) holds under [Assumption A](#) in place of no-veto power is as in footnote 22.

The proof is analogous to the proof of [Claim 1](#).

CLAIM 3. *If the equilibrium s is such that $s_i(h^t, G_t^*, \theta_t^*) = (\theta_t, f(\theta_t), v_i^f(t), 0)$ for some (θ_t^*, θ_t) , for each $i \in \mathcal{I}$, then there exists a state history θ^t and a deception π such that $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,\theta^t}^{f_\pi}(f(\theta_t), \theta_t^*)$ for all $i \in \mathcal{I}$ with $\pi_t(\theta^t, \theta_t^*) = \theta_t$ and dynamic monotonicity implies that the alternative implemented at θ_t^* is $f(\theta_t) = f(\theta_t^*)$.*

PROOF. Since $r(G_t^*; h^t) > 0$, it must be $r(G_t^*; h^t) = 1$ and the mechanism G_t^* must have been played in all periods $\tau < t$. Thus, the history h^t is uniquely determined by the strategy s and the history of realized states θ^t contained in h^t . Define $\pi_t(\theta^t, \theta_t^*) = \theta_t$ if $s_i(h^t, G_t^*, \theta_t^*) = (\theta_t, f(\theta_t), v_i^f(t), 0)$ for each $i \in \mathcal{I}$, and define $\pi_t(\theta^t, \theta_t^*) = \theta_t^*$, otherwise. Now take $\tau > t$ and consider any subsequent history h^τ of h^t (i.e., $h^\tau = (h^t, h^{t,\tau})$ for some $h^{t,\tau}$) with $q(h^\tau; s) > 0$. There are two cases. If $r(G_\tau^*; h^\tau) > 0$, define π_τ as done at t , using the history of realized states θ^τ contained in h^τ . Alternatively, if $r(G_\tau^*; h^\tau) = 0$, let $\pi_\tau(\theta^\tau, \theta_\tau) = \theta_\tau$ for all θ_τ , with θ^τ the history of realized states contained in h^τ .²⁴ Note that the constructed deception corresponds to the truth-telling deception π_τ^* whenever Rules 1.2 and 1.3 of G_τ^* apply or whenever the mechanism is P_i or D_i for some $i \in \mathcal{I}$. From [Claims 1](#) and [2](#), f is implemented whenever Rule 1.2 or 1.3 of the mechanism G_τ^* applies for any $\tau > t$. Since f is also implemented whenever the mechanism is D_i or P_i , and recalling that $f(\pi_t(\theta^t, \theta_t^*)) = f(\theta_t)$, it follows that the expected payoff of agent i under s when the state is θ_t^* at period t is $(1 - \beta_{t,T})u_i(f(\theta_t), \theta_t^*) + \beta_{t,T}v_i^f(t, \theta_t^*)$.

Now suppose that there exists $(i, x_{t,i}, v_{t,i}) \in \mathcal{I} \times X \times V_i(t)$ such that

$$(1 - \beta_{t,T})u_i(x_{t,i}, \theta_t) + \beta_{t,T}v_{t,i} \leq (1 - \beta_{t,T})u_i(f(\theta_t), \theta_t) + \beta_{t,T}v_i^f(t), \quad (3)$$

$$(1 - \beta_{t,T})u_i(x_{t,i}, \theta_t^*) + \beta_{t,T}v_{t,i} > (1 - \beta_{t,T})u_i(f(\theta_t), \theta_t^*) + \beta_{t,T}v_i^{f_\pi}(\theta^t, \theta_t^*). \quad (4)$$

If agent i deviates at (h^t, G_t^*, θ_t^*) and announces $(\theta_{t,i}, x_{t,i}, v_{t,i}, n_{t,i}) \neq (\theta_t, f(\theta_t), v_i^f(t), 0)$, then from period $t + 1$ onward he is dictatorial with probability $\lambda_i^{(t)}$ and with probability $(1 - \lambda_i^{(t)})$ the mechanism is P_i . Consequently, agent i can guarantee himself a continuation payoff of at least $v_{t,i} = \lambda_i^{(t)}\bar{v}_i + (1 - \lambda_i^{(t)})\underline{v}_i$, and thus has a profitable deviation. Therefore, for s to be an equilibrium, it must be the case that for all $(i, x_{t,i}, v_{t,i}) \in \mathcal{I} \times X \times V_i(t)$, if (3) holds, then (4) must fail. Equivalently, it must be that $L_{i,t}^f(f(\theta_t), \theta_t) \subseteq L_{i,\theta^t}^{f_\pi}(f(\theta_t), \theta_t^*)$. Since $\pi_t(\theta^t, \theta_t^*) = \theta_t$, this proves [Claim 3](#).

Since [Claims 1–3](#) are true for any period t , we conclude that s implements f . $\triangleleft$

PROOF OF THEOREM 3. We modify the canonical mechanism G_t^* as follows:

Rule 3.1. If $m_{t,i} = (\theta_t, x_t, v_t, 0)$ for all $i \in \{1, 2\}$, then $g(m) = f(\theta_t)$.

Rule 3.2a. If $m_{t,i} = (\theta_{t,i}, x_{t,i}, v_{t,i}, 0)$ and $m_{j,t} = (\theta_{t,j}, x_{t,j}, v_{t,j}, 0) \neq m_{t,i}$, then $g(m) = w$.

²⁴To see that the deception is well defined, observe that if there are two histories h^τ and $\hat{h}^\tau$ such that $q(h^\tau; r, s, p) > 0$, $q(\hat{h}^\tau; r, s, p) > 0$, $r(G_\tau^*; h^\tau) > 0$, and $r(G_\tau^*; \hat{h}^\tau) = 0$, then it must be that the history of realized states θ^τ contained in h^τ is different from the history of realized states $\hat{\theta}^\tau$ contained in $\hat{h}^\tau$ since h^τ is uniquely determined by s and θ^τ .

Rule 3.2b. If $m_{t,i} = (\theta_{t,i}, x_{t,i}, v_{t,i}, n_{t,i})$ with $n_{t,i} > 0$ and $m_{j,t} = (\theta_{t,j}, x_{t,j}, v_{t,j}, 0)$, then $g(m) = x_{t,i}$ if $(x_{t,i}, v_{t,i}) \in L_{i,t}^f(f(\theta_{t,j}), \theta_{t,j})$ and $g(m) = w$ otherwise.

Rule 3.3. If $m_{t,i} = (\theta_{t,i}, x_{t,i}, v_{t,i}, n_{t,i})$ and $m_{j,t} = (\theta_{t,j}, x_{t,j}, v_{t,j}, n_{t,j})$ with $n_{t,i} > 0$ and $n_{t,j} > 0$, then $g(m) = x_{t,i^*}$ with i^* the agent with the smallest index among the agents announcing the highest integer.

Let P^w be a mechanism that implements the outcome w , regardless of the messages. The transition rule of the dynamic mechanism regime is as follows:

- If the messages announced at period t are in Rule 3.1, the next period mechanism is the canonical mechanism with probability 1.
- If the messages announced at period t are in Rule 3.2a, then the next period mechanism is P^w with probability 1.
- If the messages announced at period t are in Rule 3.2b, then the next period mechanism is D_i (where i is the agent having announced the positive integer) with probability $\lambda_i^{(t)}$ and P^w with probability $1 - \lambda_i^{(t)}$.
- If the messages announced at period t are in Rule 3.3, then the next period mechanism is D_{i^*} with probability 1.
- If the mechanism at period t was D_i (resp., P^w), then the next period mechanism is D_i (resp., P^w) with probability 1.

As before, we compute $\lambda_i^{(t)}$ so that (i) the expected continuation payoff is $v_i^f(t)$ if $(x_{t,i}, v_{t,i}) \notin L_{i,t}^f(f(\theta_{t,j}), \theta_{t,j})$, (ii) the expected continuation payoff is $v_{t,i}$ if $(x_{t,i}, v_{t,i}) \in L_{i,t}^f(f(\theta_{t,j}), \theta_{t,j})$ and $v_{t,i} \geq v_i^w = \sum_{\theta} u_i(w, \theta)p(\theta)$, and (iii) $\lambda_{i,t} = 0$, otherwise.

To see that f is repeatedly implementable, suppose G_t^* is used at t and make the following observations:

- There are equilibrium strategies that implement f , with an equilibrium path in which G_t^* is used and all agents truthfully report $(\theta_t, f(\theta_t), v_i^f(t), 0)$ when the state is θ_t at period t . To see this, suppose that the state is $\theta_t = \theta_{t,j}$ and agent i deviates to $(\theta_{t,i}, x_{t,i}, v_{t,i}, n_{t,i})$ and either $n_{t,i} = 0$ or $(x_{t,i}, v_{t,i}) \notin L_{i,t}^f(f(\theta_{t,j}), \theta_{t,j})$. Then the alternative implemented is w and i 's continuation payoff is v_i^w . By construction, this is not a profitable deviation. If i deviates to $(\theta_{t,i}, x_{t,i}, v_{t,i}, n_{t,i})$ with $n_{t,i} > 0$ and $(x_{t,i}, v_{t,i}) \in L_{i,t}^f(f(\theta_{t,j}), \theta_{t,j})$, then the alternative adopted is $x_{t,i}$ and the continuation payoff is $v_{t,i}$; by construction, this is also not a profitable deviation (since j tells the truth).
- By condition (ii) of [Assumption A](#), any equilibrium with Rule 3.3 applying for some t must implement f .
- Any equilibrium under Rule 3.2b at $t < T$ or $t = T$ implements f . Let i be the agent reporting $n_{t,i} > 0$, let θ_t be the state reported by $j \neq i$, and let θ_t^* be the true state at t .

Consider $t < T$. First, if $\lambda_i^{(t)} < 1$, then with positive probability the outcome in all future periods is w ; agent j can profitably deviate to Rule 3.3 and guarantee himself a strictly higher continuation payoff. Second, if $\lambda_i^{(t)} = 1$, then i becomes a dictator at $t' > t$ and selects $x(\theta_{t'}) \in \max_i^{\theta_{t'}} X$ for all $t' > t$ and all $\theta_{t'}$, thus obtaining the continuation payoff $\bar{v}_i$. Write $\varphi_t(t', \theta)$ for the alternative implemented at state θ in period t' ; we have that $\varphi_t(t', \theta) \in \max_i^{\theta} X$ for all θ , for all $t' > t$. Observe that $v_i^{\varphi_t} = \bar{v}_i$. Let $(x_{i,t}, v_i(t))$ represent the second and third elements of the message sent by i at t . Agent i must have no profitable deviation, hence it must be that $\beta_{i,T} u_i(x_{i,t}, \theta_t^*) + (1 - \beta_{i,T}) \bar{v}_i \geq \beta_{i,T} u_i(y, \theta_t^*) + (1 - \beta_{i,T}) v_i$ for all $(y, v_i) \in L_{i,t}(f(\theta_t), \theta_t)$. Agent j can also deviate and become dictator himself from period t . For such a deviation not to be profitable, first it must be that $x(\theta_{t'}) \in \max_j^{\theta_{t'}} X$ for all $t' > t$ and all $\theta_{t'}$ (i.e., in all periods after t , the alternative implemented $\varphi_t(\tau, \theta)$ must therefore belong to $\max_j^{\theta} X$ for all $\tau > t$, for all θ); second, it must be that $x_{t,i} \in \max_j^{\theta_t^*} X$. The result then follows from condition (i) of [Assumption A](#).

Consider $t = T$; let the state be θ^* . It must be that $x_{T,i} \in \max_i^{\theta^*} L_i(f(\theta_{T,j}), \theta_{T,j})$. Since agent j may deviate and become dictator at $t = T$, either the deviation is profitable or by condition (ii) of [Assumption A](#), $x_{T,i} = f(\theta^*)$.

- There are no equilibria under Rule 3.2a. Let $v_i^w(t) = 0$ if $t = T$ and $v_i^w(t) = v_i^w$ otherwise. Assume that the true state at t is $\theta_t = \theta^*$ and the messages reported are $m_{t,i} = (\theta_{t,i}, x_{t,i}, v_{t,i}, 0)$ and $m_{t,j} = (\theta_{t,j}, x_{t,j}, v_{t,j}, 0)$. The alternative implemented is w and the continuation payoff vector is (v_1^w, v_2^w) . Agent i can trigger Rule 3.2b by announcing $(\theta_{t,i}, f(\theta_{t,j}), v_i^f, 1)$. Since $f(\theta_{t,j}) \in L_i(f(\theta_{t,j}), \theta_{t,j})$, it is the case that $(f(\theta_{t,j}), v_i^f) \in L_{i,t}^f(f(\theta_{t,j}), \theta_{t,j})$. Thus, the deviation yields agent i a discounted payoff of $(1 - \beta_{i,T}) u_i(f(\theta_{t,j}), \theta^*) + \beta_{i,T} v_i^f(t) > (1 - \beta_{i,T}) u_i(w, \theta^*) + \beta_{i,T} v_i^w(t)$, and hence it is profitable.
- It follows from the arguments in the proof of [Theorem 2](#) that if there are equilibria under Rule 3.1, then the dynamic lower contour sets are nested, as in the original canonical mechanism, and f is implemented. □

DEFINITION 4 (Dynamic self-selection). Let $I = 2$. For all t , all pairs $(\theta_{t,2}, \theta_{t,1})$, and all θ^t , there exists a triple $(x(\theta_{t,2}, \theta_{t,1}), v_1(\theta_{t,2}, \theta_{t,1}), v_2(\theta_{t,2}, \theta_{t,1}))$ such that $(x(\theta_{t,2}, \theta_{t,1}), v_1(\theta_{t,2}, \theta_{t,1})) \in L_{1,t}^f(f(\theta_{t,2}), \theta_{t,2})$ and $(x(\theta_{t,2}, \theta_{t,1}), v_2(\theta_{t,2}, \theta_{t,1})) \in L_{2,t}^f(f(\theta_{t,1}), \theta_{t,1})$.

Note that self-selection implies dynamic self-selection.

PROPOSITION 1. Let $I = 2$. If a social choice function f is repeatedly implementable, then it satisfies dynamic self-selection.

PROOF. Suppose that f is repeatedly implementable by the dynamic mechanism regime r . Fix an equilibrium s . Consider any period t , any history h^t , and mechanism $\langle M^{G^t}, g_t \rangle$ having positive probability of occurring on the equilibrium path, that

is, such that $q(h^t; s) > 0$ and $r(G_t; h^t) > 0$. The profile of actions $s(h^t, G_t, \hat{\theta}_t)$ must satisfy $g(s(h^t, G_t, \hat{\theta}_t)) = f(\hat{\theta}_t)$ for each $\hat{\theta}_t \in \Theta$, and the continuation payoff must be $v_t^f(t)$. This implies that

$$\begin{aligned} (1 - \beta_{t,T})u_1(f(\theta_{t,2}), \theta_{t,2}) + \beta_{t,T}v_1^f(t) \\ \geq (1 - \beta_{t,T})u_1(g(s_1(h^t, G_t, \theta_{t,1}), s_2(h^t, G_t, \theta_{t,2})), \theta_{t,2}) + \beta_{t,T}v_1(\theta_{t,2}, \theta_{t,1}), \end{aligned}$$

where $v_1(\theta_{t,2}, \theta_{t,1})$ is agent 1's continuation payoff following the deviation, and

$$\begin{aligned} (1 - \beta_{t,T})u_2(f(\theta_{t,1}), \theta_{t,1}) + \beta_{t,T}v_2^f(t) \\ \geq (1 - \beta_{t,T})u_2(g(s_1(h^t, G_t, \theta_{t,1}), s_2(h^t, G_t, \theta_{t,2})), \theta_{t,1}) + \beta_{t,T}v_2(\theta_{t,2}, \theta_{t,1}), \end{aligned}$$

where $v_2(\theta_{t,2}, \theta_{t,1})$ is agent 2's continuation payoff following the deviation. Letting $x(\theta_{t,2}, \theta_{t,1}) = g(s_1(h^t, G_t, \theta_{t,1}), s_2(h^t, G_t, \theta_{t,2}))$ completes the proof. $\square$

PROOF OF THEOREM 4. First, assume that there exists a set $\mathcal{J}$ of $I - 1$ agents such that $\bigcap_{j \in \mathcal{J}} \arg \max_{x \in X} u_j(x, \theta) \neq \emptyset$ for all θ . By no-veto power, it must be that $\{f(\theta)\} = \bigcap_{j \in \mathcal{J}} \arg \max_{x \in X} u_j(x, \theta)$ for all θ . (By no-veto power, if there exists $\{x, y\} \subseteq \bigcap_{j \in \mathcal{J}} \arg \max_{x \in X} u_j(x, \theta)$ for some θ , then $f(\theta) = x = y$, i.e., $\bigcap_{j \in \mathcal{J}} \arg \max_{x \in X} u_j(x, \theta)$ is a singleton.) The following regime implements f . At $t = 1$, all agents in $\mathcal{J}$ announce an integer and an alternative (the remaining agent is inactive). The alternative implemented at $t = 1$ is the one announced by the agent reporting the highest integer (break ties in favor of the lowest indexed agent). Moreover, the agent reporting the highest integer is dictatorial in all subsequent periods. It is routine to verify that this mechanism indeed implements f . In particular, by reporting a sufficiently large integer at $t = 1$, each agent in $\mathcal{J}$ can obtain his highest payoff with arbitrarily large probability. Therefore, it must be that the alternative implemented at each period is in $\bigcap_{j \in \mathcal{J}} \arg \max_{x \in X} u_j(x, \theta)$ for each θ , i.e., f is implemented.

Second, assume that for every set $\mathcal{J}$ of $I - 1$ agents, there exists θ and $(i, j) \in \mathcal{J} \times \mathcal{J}$ such that $\arg \max_{x \in X} u_i(x, \theta) \cap \arg \max_{x \in X} u_j(x, \theta) = \emptyset$. We make three changes to the mechanism regime adopted in the proof of Theorem 2.

First, if $T < \infty$, we replace the canonical mechanism G_T for the last period T with a slightly modified version G_T^* of the static mechanism introduced by Maskin and Sjöström (2002, p. 274). The mechanism G_T^* is $M_i = \Theta \times X \times \mathbb{N}_+ \times \{\alpha : \Theta \rightarrow X; (\alpha(\theta), 0) \in L_{i,T}^f(f(\theta), \theta)\}$, where $\mathbb{N}_+$ is the set of positive integers. There are three rules:

Rule 4.1. If $m_j = (\theta, x, 1, \cdot)$ for all $j \neq i$ and $m_i = (\theta_i, x_i, 1, \cdot)$, then $g(m) = f(\theta)$.

Rule 4.2. If $m_j = (\theta, x, 1, \cdot)$ for all $j \neq i$ and $m_i = (\theta_i, x_i, z_i, \cdot)$ with $z_i > 1$, then $g(m) = \alpha_i(\theta)$.

Rule 4.3. In all other cases, $g(m) = x_{i^*}$, where i^* is the lowest index agent among those announcing the highest integer.

Second, by no indifference for each agent i there exist θ and $\hat{y}_i$ such that $\max_x u_i(x, \theta) > u_i(\hat{y}_i, \theta)$.²⁵ We use this and modify the dictatorial mechanism D_i as

²⁵Since $p(\theta) > 0$ for all $\theta \in \Theta$, it follows that $\sum_{\theta} \max_x u_i(x, \theta) p(\theta) = \bar{v}_i > \hat{v}_i := \sum_{\theta} u_i(\hat{y}_i, \theta) p(\theta)$.

$M_i = X \times \mathbb{N}_+$, $M_j = \{\emptyset\}$ for all $j \neq i$, and $g(x_i, n_i) = (1 - 1/n_i)\mathbf{1}_{x_i} + (1/n_i)\mathbf{1}_{\hat{y}_i}$; that is, the outcome is $\hat{y}_i$ with probability $1/n_i$ and is x_i with the complementary probability.

Third, we modify the punishment mechanism P_i as $M_i = \{\emptyset\}$, $M_j = X \times \mathbb{N}_+$ for all $j \in I \setminus \{i\}$, and $g(m) = x_{j^*}$, where j^* is the lowest indexed agent that announced the highest integer. In all other aspects, the mechanism regime remains the same.

The proof that there exists an equilibrium that repeatedly implements f is essentially the same as in the proof of [Theorem 2](#); the only small change is that if G_T^* is played in the last period and the state is θ_T , then all agents report $(\theta_T, f(\theta_T), 1, \cdot)$.

To show that there are no undesirable equilibria, begin by noting that there cannot exist an equilibrium where a dictatorial regime D_i is played with positive probability on the equilibrium path; if there were, then agent i could always increase his expected payoff by announcing a higher integer, a contradiction. It follows that Rule 4.3 of G_t^* for all $t < T$ cannot be in the support of any equilibrium. In addition, for Rule 4.2 of G_t^* , $t < T$, to be in the support of an equilibrium, it must be that $\lambda_i^{(t)} = 0$ with agent i the odd man out, i.e., the mechanism transitions to P_i with probability 1 when Rule 4.2 applies. However, there exists a state θ for which the mechanism P_i has no equilibrium (since there is a pair of agents with a nonempty message space who disagree on their most preferred alternatives). It follows that any equilibrium of the game induced by the regime transitions to D_i or P_i with zero probability; that is, it corresponds to Rule 4.1 of G_t^* for all $t < T$ and thus must be in pure strategies until the last period (if $T < \infty$). This implies that if $T < \infty$, then G_T^* is played with probability 1 on the equilibrium path and hence the outcome at T must correspond to a (mixed strategy) Nash equilibrium of G_T^* .

As argued by [Maskin and Sjöström \(2002\)](#), in G_T^* an “agent i has nothing to loose from setting” (i) $\alpha_i(\theta)$ equal to his favorite outcome in the lower contour set of $f(\theta)$ at θ , (ii) $x_i \in \arg\max_{x \in X} u_i(x, \theta_T^*)$, where θ_T^* is the true state, and (iii) “ z_i larger than any integer announced with positive probability by any other agent.” First, this implies that when Rule 4.2 or Rule 4.3 of G_T^* applies at state θ_T^* , then by no-veto power it must be that the alternative implemented is $f(\theta_T^*)$. Second, it implies that when Rule 4.1 applies at θ_T^* , it must be that the state reported by $I - 1$ agents is the same, denote it by $\pi_T(\cdot, \theta_T^*)$, the alternative implemented is $f(\pi_T(\cdot, \theta_T^*))$, and there is no alternative in $L_{i,T}^f(f(\pi_T(\cdot, \theta_T^*)), \pi_T(\cdot, \theta_T^*))$ that is preferred by agent i to $f(\pi_T(\cdot, \theta_T^*))$; that is, $L_{i,T}^f(f(\pi_T(\cdot, \theta_T^*)), \pi_T(\cdot, \theta_T^*)) \subseteq L_{i,T}^f(f(\pi_T(\cdot, \theta_T^*)), \theta_T^*)$.²⁶

Now consider any equilibrium σ in behavioral strategies. From the above argument, the mechanism adopted is G_t^* at any t . In addition, at all $t < T$, histories h^t , and states θ_t , the mixed action $\sigma(h^t, G_t^*, \theta_t)$ is pure and corresponds to Rule 4.1 of G_t^* . It follows that we can associate with any state profile θ^t a unique public history $h_D^t(\theta^t)$ over mechanisms, messages reported, and alternatives implemented. We now define the deceptions induced by σ .

For any $t < T$, for any (θ^t, θ_t) , we simply define the map $\pi_t(\theta^t, \theta_t) = \theta'_t$, where θ'_t is the common state reported by at least $I - 1$ agents at the history $(h_D^t(\theta^t), \theta^t, G_t^*, \theta_t)$ under σ . For any (θ^T, θ_T) , we define a distribution $q(\theta^T, \theta_T) \in \Delta(\Theta)$ over maps

²⁶Hence, if we assumed Maskin monotonicity, f would also be implemented at T when Rule 4.1 applies, but dynamic monotonicity does not imply Maskin monotonicity, as shown by [Example 1](#).

$\pi_T(\theta^T, \theta_T) \in \Theta$ such that (i) $\pi_T(\theta^T, \theta_T) = \theta_T$ with probability $q(\theta^T, \theta_T)[\theta_T]$ given by the sum of $\sigma(h_D^T(\theta^T), \theta^T, G_T^*, \theta_T)[m]$ over all messages m such that either Rule 4.2 or 4.3 applies or Rule 4.1 applies with θ_T being the common state reported by at least $I - 1$ agents, and (ii) $\pi_T(\theta^T, \theta_T) = \theta'_T$ with probability $q(\theta^T, \theta_T)[\theta'_T]$ given by the sum of $\sigma(h_D^T(\theta^T), \theta^T, G_T^*, \theta_T)[m]$ over all messages m such that Rule 4.1 applies with θ'_T being the common state reported by at least $I - 1$ agents, for all $\theta'_T \neq \theta_T$.

We have thus defined a distribution over a set of dynamic deceptions, where deception π^k , $k \in \mathcal{K}$, has probability q^k . For instance, if π^k is such that $\pi_T^k(\theta^T, \theta_T) = \theta_T$ for all (θ^T, θ_T) , the probability q^k is $\times_{(\theta^T, \theta_T)} q(\theta^T, \theta_T)[\theta_T]$.

The expected payoff of agent i at $t < T$, when the state is θ_t^* and i selects the pure strategy s_i in the support of σ_i while all other agents follow their behavioral strategies σ_{-i} (and hence report state $\pi_t(\theta^t, \theta_t^*)$), is $(1 - \beta_{t,T})u_i(f(\pi_t(\theta^t, \theta_t^*)), \theta_t^*) + \beta_{t,T} \sum_{k \in \mathcal{K}} q^k v_i^{f, \pi^k}(\theta^t, \theta_t^*)$. Suppose there exists $(i, x_{t,i}, v_{t,i}) \in \mathcal{I} \times X \times V_i(t)$ such that

$$(1 - \beta_{t,T})u_i(x_{t,i}, \pi_t(\theta^t, \theta_t^*)) + \beta_{t,T}v_{t,i} \leq (1 - \beta_{t,T})u_i(f(\pi_t(\theta^t, \theta_t^*)), \pi_t(\theta^t, \theta_t^*)) + \beta_{t,T}v_i^f(t) \quad (5)$$

$$(1 - \beta_{t,T})u_i(x_{t,i}, \theta_t^*) + \beta_{t,T}v_{t,i} > (1 - \beta_{t,T})u_i(f(\pi_t(\theta^t, \theta_t^*)), \theta_t^*) + \beta_{t,T} \sum_{k \in \mathcal{K}} q^k v_i^{f, \pi^k}(\theta^t, \theta_t^*). \quad (6)$$

If agent i deviates at (h^t, G_t^*, θ_t^*) , $t < T$, and sends the message $(\theta_{t,i}, x_{t,i}, v_{t,i}, n_{t,i}) \neq (\pi_t(\theta^t, \theta_t^*), f(\pi_t(\theta^t, \theta_t^*)), v_i^f(t), 0)$, then from period $t + 1$ onward he is dictatorial with probability $\lambda_i^{(t)}$ and with probability $(1 - \lambda_i^{(t)})$ the mechanism is P_i . By selecting an arbitrarily large integer when the mechanism is D_i , agent i obtains at least a continuation payoff arbitrarily close to $v_{t,i} = \lambda_i^{(t)}\bar{v}_i + (1 - \lambda_i^{(t)})\underline{v}_i$, and thus has a profitable deviation. Hence, for σ to be an equilibrium, it must be that for all $(i, x_{t,i}, v_{t,i}) \in \mathcal{I} \times X \times V_i(t)$, if (5) holds, then (6) fails; that is, it must be that $L_{i,t}^f(f(\pi_t(\theta^t, \theta_t^*)), \pi_t(\theta^t, \theta_t^*)) \subseteq L_{i,\theta^t}^{f, \pi^k}(f(\pi_t(\theta^t, \theta_t^*)), \theta_t^*)$ for at least one deception π^k , for all $t < T$. We have already established that it must also be $L_{i,T}^f(f(\pi_T(\theta^T, \theta_T^*)), \pi_T(\theta^T, \theta_T^*)) \subseteq L_{i,T}^f(f(\pi(\theta^T, \theta_T^*)), \theta_T^*) = L_{i,\theta^T}^{f, \pi^k}(f(\pi_T(\theta^T, \theta_T^*)), \theta_T^*)$. Dynamic monotonicity then implies $f(\pi_t(\theta^t, \theta_t^*)) = f(\theta_t^*)$ for all $t \in \mathcal{T}$. This concludes the proof of the theorem. $\square$

REFERENCES

- Abreu, Dilip and Arunava Sen (1990), “Subgame perfect implementation: A necessary and almost sufficient condition.” *Journal of Economic Theory*, 50, 285–299. [266]
- Aghion, Philippe, Drew Fudenberg, Richard Holden, Takashi Kunimoto, and Olivier Tercieux (2012), “Subgame-perfect implementation under information perturbations.” *Quarterly Journal of Economics*, 127, 1843–1881. [254, 267]

- Åzakis, Helmut and Péter Vida (2015), “Repeated implementation.” Unpublished paper, University of Mannheim. [263]
- Benoît, Jean-Pierre and Vijay Krishna (1987), “Nash equilibria of finitely repeated games.” *International Journal of Game Theory*, 16, 197–204. [257]
- Bergemann, Dirk and Maher Said (2011), “Dynamic auctions.” In *Wiley Encyclopedia of Operations Research and Management Science* (James J. Cochran, ed.), 1511–1522, John Wiley & Sons, New York. [250]
- Chambers, Christopher P. (2004), “Virtual repeated implementation.” *Economics Letters*, 83, 263–268. [251]
- Dutta, Bhaskar and Arunava Sen (1991), “A necessary and sufficient condition for two-person Nash implementation.” *Review of Economic Studies*, 58, 121–128. [264]
- González-Díaz, Julio (2006), “Finitely repeated games: A generalized Nash folk theorem.” *Games and Economic Behavior*, 55, 100–111. [257]
- Jackson, Matthew O. (2001), “A crash course in implementation theory.” *Social Choice and Welfare*, 18, 655–708. [250]
- Kalai, Ehud and John O. Ledyard (1998), “Repeated implementation.” *Journal of Economic Theory*, 83, 308–317. [251]
- Lee, Jihong and Hamid Sabourian (2011), “Efficient repeated implementation.” *Econometrica*, 79, 1967–1994. [250, 251, 254, 256, 257, 260]
- Lee, Jihong and Hamid Sabourian (2013), “Repeated implementation with incomplete information.” Unpublished paper, Cambridge University and Seoul University. [251]
- Mailath, George J. and Larry Samuelson (2006), *Repeated Games and Reputations: Long-run Relationships*. Oxford University Press, Oxford. [271]
- Maskin, Eric (1999), “Nash equilibrium and welfare optimality.” *Review of Economic Studies*, 66, 23–38. [250, 262, 263]
- Maskin, Eric and Tomas Sjöström (2002), “Implementation theory.” In *Handbook of Social Choice and Welfare* (Kenneth J. Arrow, Amartya Sen, and Kotaro Suzumura, eds.), 237–288, North Holland, Amsterdam. [250, 266, 281, 282]
- Maskin, Eric and Jean Tirole (1999), “Unforeseen contingencies and incomplete contracts.” *Review of Economic Studies*, 66, 83–114. [267]
- Mezzetti, Claudio and Ludovic Renou (2012), “Implementation in mixed Nash equilibrium.” *Journal of Economic Theory*, 147, 2357–2375. [263]
- Moore, John H. and Rafael Repullo (1988), “Subgame perfect implementation.” *Econometrica*, 56, 1191–1220. [264, 266]
- Renou, Ludovic and Tristan Tomala (2015), “Approximate implementation in Markovian environments.” *Journal of Economic Theory*, 159, 401–442. [251]

Serrano, Roberto (2004), “The theory of implementation of social choice rules.” *SIAM Review*, 46, 377–414. [250]

Co-editor George J. Mailath handled this manuscript.

Manuscript received 7 October, 2014; final version accepted 26 January, 2016; available online 6 February, 2016.