

Efficient Repeated Implementation

Jihong Lee*

Hamid Sabourian[†]

Yonsei University and
Birkbeck College, London

University of Cambridge

September 2009

Abstract

This paper examines repeated implementation of a social choice function (SCF) with infinitely-lived agents whose preferences are determined randomly in each period. An SCF is repeated-implementable in (Bayesian) Nash equilibrium if there exists a sequence of (possibly history-dependent) mechanisms such that (i) its equilibrium set is non-empty and (ii) every equilibrium outcome corresponds to the desired social choice at every possible history of past play and realizations of uncertainty. We first show, with minor qualifications, that in the complete information environment an SCF is repeated-implementable if and only if it is efficient. We then extend this result to the incomplete information setup. In particular, it is shown that in this case efficiency is sufficient to ensure the characterization part of repeated implementation. For the existence part, incentive compatibility is sufficient but not necessary. In the case of interdependent values, existence can also be established with an intuitive condition stipulating that deviations can be detected by at least one agent other than the deviator. Our incomplete information analysis can be extended to incorporate the notion of ex post equilibrium.

JEL Classification: A13, C72, C73, D78

Keywords: Repeated implementation, Nash implementation, Bayesian implementation, Ex post implementation, Efficiency, Incentive compatibility, Identifiability

*School of Economics, Yonsei University, Seoul 120-749, Korea, Jihong.Lee@yonsei.ac.kr

[†]Faculty of Economics, Cambridge, CB3 9DD, United Kingdom, Hamid.Sabourian@econ.cam.ac.uk

1 Introduction

Implementation theory, sometimes referred to as the theory of *full* implementation, has been concerned with designing mechanisms, or game forms, that implement desired social choices in *every* equilibrium of the mechanism. Numerous characterizations of implementable social choice rules have been obtained in one-shot settings in which agents interact only once. However, many real world institutions, from voting and markets to contracts, are used repeatedly by their participants, and implementation theory has yet to offer much to the question of what is generally implementable in repeated contexts (see, for example, Jackson [11]).¹

This paper examines repeated implementation in environments in which agents' preferences are determined stochastically across time and, therefore, a sequence of mechanisms need to be devised in order to repeatedly implement desired social choices. In our setup, the agents are infinitely-lived and their preferences represented by state-dependent utilities with the state being drawn randomly in each period from an identical prior distribution. Utilities are not necessarily transferable. The information structure of our setup is general and allows for both private and interdependent values as well as correlated types, including the case of complete information.

In the one-shot implementation problem the critical conditions for implementing a social choice rule are (Maskin) monotonicity for the complete information case and Bayesian monotonicity (an extension of Maskin monotonicity) and incentive compatibility for the incomplete information case. These conditions are necessary for implementation and, together with some minor assumptions, are also sufficient. However, they can be very strong restrictions as many desirable social choice rules fail to satisfy them (see the surveys of Jackson [11], Maskin and Sjöström [20] and Serrano [28], among others).

As is the case between one-shot and repeated games a repeated implementation problem introduces fundamental differences to what we have learned about implementation in the one-shot context. In particular, one-shot implementability does not imply repeated implementability if the agents can co-ordinate on histories, thereby creating other, possibly unwanted, equilibria.

¹The literature on dynamic mechanism design does not address the issue of full-implementation since it is concerned only with establishing the existence of a single equilibrium of some mechanism that possesses the desired properties.

To gain some intuition, consider a social choice function that satisfies sufficiency conditions for Nash implementation in the one-shot complete information setup (e.g. monotonicity and no veto power) and a mechanism that implements it (e.g. the mechanism proposed by Maskin [18]). Suppose now that the agents play this mechanism repeatedly. Assume also that in each period a state is drawn independently from a fixed distribution and the realizations are complete information.² Then, the game played by the agent is simply a repeated game with random states. Since in the stage game every Nash equilibrium outcome corresponds to the desired outcome in each state, this repeated game has an equilibrium in which each agent plays the desired action at each period/state regardless of past histories. However, we also know from the study of repeated games (see Mailath and Samuelson [16]) that unless minmax expected utility profile of the stage game lies on the efficient payoff frontier of the repeated game, by the Folk theorem, there will be many equilibrium paths along which unwanted outcomes are implemented. Thus, the conditions that guarantee one-shot implementation are not sufficient for repeated implementation. Our results below also show that they are not necessary either.

Given the multiple equilibria and collusion possibilities in repeated environments, at first glance, implementation in such settings seems a very daunting task. But our understanding of repeated interactions also provides us with several clues as to how such repeated implementation may be achieved. First, a critical condition for repeated implementation is likely to be some form of efficiency of the social choices, that is, the payoff profile of the social choice function ought to lie on the efficient frontier of the corresponding repeated game/implementation payoffs. Second, we need to devise a sequence of mechanisms such that, roughly speaking, the agents' individually rational payoffs also coincide with the efficient payoff profile of the social choice function.

While repeated play introduces the possibility of co-ordinating on histories by the agents, thereby creating difficulties towards full repeated implementation, it also allows for more structure in the mechanisms that the planner can enforce. We introduce a sequence of mechanisms, or a *regime*, such that the mechanism played in a given period depends on the past history of mechanisms played and the agents' corresponding actions. This way the infinite future gives the planner additional leverage: the planner can alter the future mechanisms in a way that rewards desirable behavior while punishing the undesirable.

²A detailed example is provided in Section 3 below.

Formally, we consider repeated implementation of a social choice function (henceforth, called SCF) in the following sense: there exists a regime such that (i) its equilibrium set is non-empty and (ii) in any equilibrium of the regime, the desired social choice is implemented at every possible history of past play of the regime and realizations of states. A weaker notion of repeated implementation seeks the equilibrium continuation payoff (discounted average expected utilities) of each agent at every possible history to correspond precisely to the one-shot payoff (expected utility) of the social choices. Our complete information analysis adopts Nash equilibrium as the solution concept; the incomplete information analysis considers repeated implementation in Bayesian Nash equilibrium.³

Our results establish for the general, complete or incomplete information, environment that, with some minor qualifications, the characterization part of repeated implementation is achievable if and only if the SCF is efficient. Our main messages are most cleanly delivered in the complete information setup. Here, we first demonstrate the following necessity result: if the agents are sufficiently patient and an SCF is repeated-implementable, it cannot be Pareto-dominated (in terms of expected utilities) by another SCF whose *range* belongs to that of the desired SCF. Just as the theory of repeated game suggests, the agents can indeed “collude” in our repeated implementation setup if there is a possibility of collective benefits.

We then present the paper’s main result for the complete information case: under some minor conditions, every efficient SCF can be repeatedly implemented. This sufficiency result is obtained by constructing for each SCF a canonical regime in which, at any history along an equilibrium path, each agent’s continuation payoff has a lower bound equal to his payoff from the SCF, thereby ensuring the individually rational payoff profile in any continuation game to be no less than the desired profile. It then follows that if the latter is located on the efficient frontier the agents cannot sustain any collusion away from the desired payoff profile; moreover, if there is a unique SCF associated with such payoffs then repeated implementation of desired outcomes is achieved.

The construction of the canonical regime involves two steps. We first show, for each player i , that there exists a regime S^i in which the player obtains a payoff exactly equal to that from the SCF and, then, embed this into the canonical regime such that each

³Thus, our solution concepts do not rely on imposing credibility and/or particular belief specification *off-the-equilibrium*, which have been adopted elsewhere to sharpen predictions (for example, Moore and Repullo [23], Abreu and Sen [2] and Bergin and Sen [6]).

agent i can always induce S^i in the continuation game by an appropriate deviation from his equilibrium strategy. The first step is obtained by applying Sorin’s [27] observation that with infinite horizon any payoff can be generated exactly by the discounted average payoff from some sequence of outcomes, as long as the discount factor is sufficiently large.⁴ The second step is obtained by allowing each agent the possibility of making himself the “odd-one-out”, thereby inducing S^i in the continuation play, in any equilibrium.

These arguments also enable us to handle the incomplete information case, delivering results that closely parallel those above. These analogous results are obtained for very general incomplete information setup in which no restrictions are imposed on the information structure (both private value and interdependent value cases are allowed, as well as correlated types). Also, as with complete information, no (Bayesian) monotonicity assumption is needed for this result. This is important to note because monotonicity in the incomplete information setup is a very demanding restriction (see Serrano [28]).

With incomplete information, however, there are several additional issues. First, we evaluate repeated implementation in terms of *expected* continuation payoffs computed at the beginning of a regime. This is because continuation payoffs in general depend on an agent’s ex post beliefs about the others’ past private information at different histories but we do not want our solution concept to depend on such beliefs.

Second, although efficiency pins down payoffs of every equilibrium, it still remains to establish existence of an equilibrium in the canonical regime. Incentive compatibility offers one natural sufficiency condition for existence. However, we also demonstrate how we can do without incentive compatibility. In the interdependent value case, where there can be a serious conflict between efficiency and incentive compatibility (see Maskin [17] and Jehiel and Moldovanu [14]), the same set of results are obtained by replacing incentive compatibility with an intuitive condition that we call *identifiability*, if the agents are sufficiently patient. Identifiability requires that when agents announce types and all but one agent reports his type truthfully then, upon learning his utility at the end of the period, someone other than the untruthful odd-one-out will discover that there was a lie. Given identifiability, we construct another regime that, while maintaining the desired payoff properties of its equilibrium set, admits a truth-telling equilibrium based on incentives of repeated play instead of one-shot incentive compatibility of the SCF. Such

⁴In our setup, the required threshold on discount factor is one half and, therefore, our main sufficiency results do not in fact depend on an arbitrarily large discount factor.

incentives involve punishment when someone misreports his type.

Third, we can extend our incomplete information analysis by adopting the notion of *ex post* equilibrium, thereby requiring the agents' strategies to be mutually optimal for every possible realization of past and present types and not just for some given distribution (see Bergemann and Morris [4][5] for some recent contributions to one-shot implementation that address this issue). The precise nature of the agents' private information are not relevant in our constructive arguments and, in fact, the more stringent restrictions imposed by *ex post* equilibrium enable us to derive a sharper set of results that are closer to the results with complete information.

To this date, only few papers address the problem of repeated implementation. Kalai and Ledyard [15] and Chambers [7] ask the question of implementing an infinite sequence of outcomes when the agents' preferences are fixed. Kalai and Ledyard [15] find that, if the social planner is more patient than the agents and, moreover, is interested only in the long-run implementation of a sequence of outcomes, he can elicit the agents' preferences truthfully in dominant strategies. Chambers [7] applies the intuitions behind the virtual implementation literature to demonstrate that, in a continuous time, complete information setup, any outcome sequence that realizes every feasible outcome for a positive amount of time satisfies monotonicity and no veto power and, hence, is Nash implementable.

In these models, however, there is only one piece of information to be extracted from the agents who, therefore, do not interact repeatedly themselves. More recently, Jackson and Sonnenschein [12] consider linking a specific, *independent private values*, Bayesian implementation problem with a large, but finite, number of independent copies of itself. If the linkage takes place through time, their setup can be interpreted as a particular finitely repeated implementation problem. The authors restrict their attention to a sequence of revelation mechanisms in which each agent is budgeted in his choice of messages according to the prior distribution over his possible types. They find that for any *ex ante* Pareto efficient SCF all equilibrium payoffs of such a budgeted mechanism must approximate the target payoff profile corresponding to the SCF, as long as the agents are sufficiently patient and the horizon sufficiently long.

In contrast to Jackson and Sonnenschein [12], our setup deals with infinitely-lived agents and a fully general information structure that allows for interdependent values as well as complete information. In terms of the results, we derive precise, rather than approximate, repeated implementation of an efficient SCF at every possible history of

the regime, not just in terms of payoffs computed at the outset. Our main results do not require the discount factor to be arbitrarily large. Furthermore, these results are obtained with arguments that are very much distinct from those of [12].

The paper is organized as follows. Section 2 first introduces one-shot implementation with complete information which will provide the basic definitions and notation throughout the paper. Section 3 then describes the infinitely repeated implementation problem and presents our main results for the complete information setup. In Section 4, we extend the analysis to incorporate incomplete information. Section 5 offers some concluding remarks about potential extensions of our analysis. Some proofs are relegated to an Appendix. Also, we provide a Supplementary Material to present some results and proofs whose details are left out for space reasons.

2 Complete information: preliminaries

Let I be a finite, non-singleton set of agents; with some abuse of notation, I also denotes the cardinality of this set. Let A be a finite set of outcomes and Θ be a finite, non-singleton set of the possible states and p denote a probability distribution defined on Θ such that $p(\theta) > 0$ for all $\theta \in \Theta$. Agent i 's state-dependent utility function is given by $u_i : A \times \Theta \rightarrow \mathbb{R}$. An *implementation problem*, $\mathcal{P}$, is a collection $\mathcal{P} = [I, A, \Theta, p, (u_i)_{i \in I}]$.

An SCF f in an implementation problem $\mathcal{P}$ is a mapping $f : \Theta \rightarrow A$ such that $f(\theta) \in A$ for any $\theta \in \Theta$. The *range* of f is the set $f(\Theta) = \{a \in A : a \in f(\theta) \text{ for some } \theta \in \Theta\}$. Let F denote the set of all possible SCFs and, for any $f \in F$, define $F(f) = \{f' \in F : f'(\Theta) \subseteq f(\Theta)\}$ as the set of all SCFs whose range belongs to $f(\Theta)$.

For an outcome $a \in A$, define $v_i(a) = \sum_{\theta \in \Theta} p(\theta) u_i(a, \theta)$ as its expected utility, or (one-shot) payoff, to agent i . Similarly, though with some abuse of notation, for an SCF f define $v_i(f) = \sum_{\theta \in \Theta} p(\theta) u_i(f(\theta), \theta)$. Denote the profile of payoffs associated with f by $v(f) = (v_i(f))_{i \in I}$. Let $V = \{v(f) \in \mathbb{R}^I : f \in F\}$ be the set of expected utility profiles of all possible SCFs. Also, for a given $f \in F$, let $V(f) = \{(v_i(f'))_{i \in I} \in \mathbb{R}^I : f' \in F(f)\}$ be the set of payoff profiles of all SCFs whose ranges belong to the range of f . We write $co(V)$ and $co(V(f))$ for the convex hulls of the two sets, respectively.

A payoff profile $v' = (v'_1, \dots, v'_I) \in co(V)$ is said to Pareto dominate another profile $v = (v_1, \dots, v_I)$ if $v'_i \geq v_i$ for all i with the inequality being strict for at least one agent. Furthermore, v' *strictly* Pareto dominates v if the inequality is strict for *all* i . An *efficient*

SCF is defined as follows.

Definition 1 *An SCF f is efficient if there exists no $v' \in co(V)$ that Pareto dominates $v(f)$; f is strictly efficient if it is efficient and there exists no $f' \in F$, $f' \neq f$, such that $v(f') = v(f)$.*

Our notion of efficiency is similar to *ex ante* Pareto efficiency used by Jackson and Sonnenschein [12]. The difference is that we define efficiency over the *convex hull* of the set of expected utility profiles of all possible SCFs. As will shortly become clear, this reflects the set of payoffs that can be obtained in an infinitely repeated implementation problem (i.e. discounted average expected utility profiles).⁵

We also define efficiency *in the range* as follows.

Definition 2 *An SCF f is efficient in the range if there exists no $v' \in co(V(f))$ that Pareto dominates $v(f)$; f is strictly efficient on the range if it is efficient in the range and there exists no $f' \in F(f)$, $f' \neq f$, such that $v(f') = v(f)$.*

As a benchmark, we next specify Nash implementation in the one-shot context. A mechanism is defined as $g = (M^g, \psi^g)$, where $M^g = M_1^g \times \cdots \times M_I^g$ is a cross product of message spaces and $\psi^g : M^g \rightarrow A$ is an outcome function such that $\psi^g(m) \in A$ for any message profile $m = (m_1, \dots, m_I) \in M^g$. Let G be the set of all feasible mechanisms.

Given a mechanism $g = (M^g, \psi^g)$, we denote by $\mathcal{N}_g(\theta) \subseteq M^g$ the set of Nash equilibria of the game induced by g in state θ . We then say that an SCF f is *Nash implementable* if there exists a mechanism g such that, for all $\theta \in \Theta$, $\psi^g(m) = f(\theta)$ for all $m \in \mathcal{N}_g(\theta)$.

The seminal result on Nash implementation is due to Maskin [18]: (i) If an SCF f is Nash implementable, f satisfies monotonicity; (ii) If $I \geq 3$, and if f satisfies monotonicity and no veto power, f is Nash implementable.⁶

Monotonicity can be a strong concept.⁷ It may not even be consistent with efficiency in standard problems such as voting or auction, and as result, efficient SCFs may not be

⁵Clearly an efficient f is ex post Pareto efficient in that, given state θ , $f(\theta)$ is Pareto efficient. An ex post Pareto efficient SCF needs not however be efficient.

⁶An SCF f is *monotonic* if, for any $\theta, \theta' \in \Theta$ and $a = f(\theta)$ such that $a \neq f(\theta')$, there exist some $i \in I$ and $b \in A$ such that $u_i(a, \theta) \geq u_i(b, \theta)$ and $u_i(a, \theta') < u_i(b, \theta')$. An SCF f satisfies *no veto power* if, whenever i, θ and a are such that $u_j(a, \theta) \geq u_j(b, \theta)$ for all $j \neq i$ and all $b \in A$, then $a = f(\theta)$.

⁷Some formal results showing the restrictiveness of monotonicity can be found in Mueller and Satterthwaite [25], Dasgupta, Hammond and Maskin [9] and Saijo [26].

implementable in such one-shot settings. To illustrate this, consider an implementation problem where $I = \{1, 2, 3, 4\}$, $A = \{a, b, c\}$, $\Theta = \{\theta', \theta''\}$ and the agents' state-contingent utilities are given below:

	θ'				θ''			
	$i = 1$	$i = 2$	$i = 3$	$i = 4$	$i = 1$	$i = 2$	$i = 3$	$i = 4$
a	3	2	1	3	3	2	1	3
b	1	3	2	2	2	3	3	2
c	2	1	3	1	1	1	2	1

The SCF f is such that $f(\theta') = a$ and $f(\theta'') = b$. Notice that f is *utilitarian* (i.e. maximizes the sum of agents' utilities) and, hence, (strictly) efficient; moreover, in a voting context, such social objectives can be interpreted as representing a scoring rule, such as the Borda count. However, the SCF is not monotonic. The position of outcome a does not change in any agent's preference ordering across the two states and yet a is $f(\theta')$ but not $f(\theta'')$.

Consider another example with three players/outcomes and the following utilities:

	θ'			θ''		
	$i = 1$	$i = 2$	$i = 3$	$i = 1$	$i = 2$	$i = 3$
a	30	0	0	10	0	0
b	0	10	0	0	30	0
c	0	0	20	0	0	20

This is an auction without transfers. Outcomes a , b and c represent the object being awarded to agent 1, 2 and 3, respectively, and each agent derives positive utility if and only if he obtains the object. In this case, the relative ranking of outcomes does not change for any agent but the social choice may vary with agents' preference intensity such that $f(\theta') = a$ and $f(\theta'') = b$. Here, such an SCF, which is clearly efficient, has no hope of satisfying monotonicity, or even *ordinality*, which allows for virtual implementation (Matsushima [21] and Abreu and Sen [3]).⁸

⁸An SCF f is ordinal if, whenever $f(\theta) \neq f(\theta')$, there exist some individual i and two outcomes (lotteries) $a, b \in A$ such that $u_i(a, \theta) \geq u_i(b, \theta)$ and $u_i(a, \theta') < u_i(b, \theta')$.

3 Complete information: repeated implementation

3.1 A motivating example

We begin by discussing an illustrative example. Consider the following case: $I = \{1, 2, 3\}$, $A = \{a, b, c\}$, $\Theta = \{\theta', \theta''\}$ and the agents' state-contingent utilities are given below:

	θ'			θ''		
	$i = 1$	$i = 2$	$i = 3$	$i = 1$	$i = 2$	$i = 3$
a	4	2	2	3	1	2
b	0	3	3	0	4	4
c	0	0	4	0	2	3

The SCF f is such that $f(\theta') = a$ and $f(\theta'') = b$. This SCF is efficient, monotonic and satisfies no veto power. The Maskin mechanism, $\mathcal{M} = (M, \psi)$, for f is defined as follows: $M_i = \Theta \times A \times \mathbb{Z}_+$ (where $\mathbb{Z}_+$ is the set of non-negative integers) for all i and ψ satisfies

1. if $m_i = (\theta, f(\theta), 0)$ for all i , $\psi(m) = f(\theta)$;
2. if there exists some i such that $m_j = (\theta, f(\theta), 0)$ for all $j \neq i$ and $m_i = (\cdot, \tilde{a}, \cdot) \neq m_j$, $\psi(m) = \tilde{a}$ if $u_i(f(\theta), \theta) \geq u_i(\tilde{a}, \theta)$ and $\psi(m) = f(\theta)$ if $u_i(f(\theta), \theta) < u_i(\tilde{a}, \theta)$;
3. if $m = ((\theta^i, a^i, z^i))_{i \in I}$ is of any other type and i is lowest-indexed agent among those who announce the highest integer, $\psi(m) = a^i$.

By monotonicity and no veto power of f , for each θ , the unique Nash equilibrium of M consists of each agent announcing $(\theta, f(\theta), 0)$, thereby inducing outcome $f(\theta)$.

Next, consider the infinitely repeated version of Maskin mechanism, where in each period state θ is drawn randomly and the agents play the same Maskin mechanism. Then, the players face an infinitely repeated game with random state at each period. Clearly, this repeated game admits an equilibrium in which the agents play the unique Nash equilibrium of the stage game in each state regardless of past history, thereby implementing f in each period. However, if the agents are sufficiently patient, there will be other equilibria and the SCF cannot be (uniquely) implemented.

For instance, consider the following repeated game strategies which implement outcome b in both states of each period. Each agent reports $(\theta'', b, 0)$ in each state/period

with the following punishment schemes: (i) if either agent 1 or 2 deviates then each agent ignores the deviation and continues to report the same; (ii) if agent 3 deviates then each agent plays the stage game Nash equilibrium in each state/period thereafter independently of subsequent history.

It is easy to see that neither agent 1 nor agent 2 has an incentive to deviate: although agent 1 would prefer a over b in both states, the rules of $\mathcal{M}$ do not allow implementation of a from his unilateral deviation; on the other hand, agent 2 is getting his most preferred outcome in each state. If the discount factor is sufficiently large, agent 3 does not want to deviate either. He can indeed alter the implemented outcome in state θ' and obtain c instead of b (after all, this is the agent whose preference reversal supports monotonicity). However, such a deviation would be met by (credible) punishment in which his continuation payoff is a convex combination of 2 (in θ') and 4 (in θ''), which is less than the equilibrium continuation payoff.

In the above example, we have deliberately chosen an SCF that is efficient (as well as monotonic and satisfies no veto power). As a result, the Maskin mechanism in the one-shot framework induces a unique Nash equilibrium payoffs on its efficient frontier. Despite this, in the repeated framework, there are many equilibria and the SCF cannot be implemented with a repeated version of Maskin mechanism. The reason for this lack of implementation in this example is that, in the Maskin game form, the Nash equilibrium payoffs are different from the minmax payoffs. For instance, agent 1's minmax utility in θ' is equal to 0, resulting from $m_2 = m_3 = (\theta'', f(\theta''), 0)$, which is less than his utility from $f(\theta') = a$; in θ'' , minmax utilities of agents 2 and 3, which both equal 2, are below their respective utilities from $f(\theta'') = b$. As a result, the set of individually rational payoffs in the repeated setup is not singleton, and one can obtain numerous equilibrium paths/payoffs with sufficiently patient agents.

The above example highlights the fundamental difference between repeated and one-shot implementation, and suggests that one-shot implementability, characterized by monotonicity and no veto power of an SCF, may be irrelevant for repeated implementability. Our understanding of repeated interactions and the multiplicity of equilibria gives us two clues. First, a critical condition for repeated implementation is likely to be some form of efficiency of the social choices, that is, the payoff profile of the SCF ought to lie on the efficient frontier of the repeated game/implementation payoffs. Second, we want to devise a sequence of mechanisms such that, roughly speaking, the agents' individually rational

payoffs also coincide with the efficient payoff profile of the SCF. In what follows, we shall demonstrate that these intuitions are indeed correct and, moreover, achievable.

3.2 Definitions

An *infinitely repeated implementation problem* is denoted by $\mathcal{P}^\infty$, representing infinite repetitions of the implementation problem $\mathcal{P} = [I, A, \Theta, p, (u_i)_{i \in I}]$. Periods are indexed by $t \in \mathbb{Z}_{++}$. In each period, the state is drawn from Θ from an independent and identical probability distribution p .

An (uncertain) infinite sequence of outcomes is denoted by $a^\infty = (a^{t,\theta})_{t \in \mathbb{Z}_{++}, \theta \in \Theta}$, where $a^{t,\theta} \in A$ is the outcome implemented in period t and state θ . Let A^∞ denote the set of all such sequences. Agents' preferences over alternative infinite sequences of outcomes are represented by discounted average expected utilities. Formally, $\delta \in (0, 1)$ is the agents' common discount factor, and agent i 's (repeated game) payoffs are given by a mapping $\pi_i : A^\infty \rightarrow \mathbb{R}$ such that

$$\pi_i(a^\infty) = (1 - \delta) \sum_{t \in \mathbb{Z}_{++}} \sum_{\theta \in \Theta} \delta^{t-1} p(\theta) u_i(a^{t,\theta}, \theta).$$

It is assumed that the structure of an infinitely repeated implementation problem (including the discount factor) is common knowledge among the agents and, if there is one, the social planner. The realized state in each period is complete information among the agents but unobservable to an outsider.

We want to repeatedly implement an SCF in each period by devising a mechanism for each period. A *regime* specifies a sequence of mechanisms contingent on the publicly observable history of mechanisms played and the agents' corresponding actions. It is assumed that a planner, or the agents themselves, can commit to a regime at the outset.

Some notation is needed to formally define a regime. Given a mechanism $g = (M^g, \psi^g)$, define $\mathcal{E}^g \equiv \{(g, m)\}_{m \in M^g}$, and let $\mathcal{E} = \cup_{g \in G} \mathcal{E}^g$. Let $H^t = \mathcal{E}^{t-1}$ (the $(t-1)$ -fold Cartesian product of $\mathcal{E}$) represent the set of all possible histories of mechanisms played and the agents' corresponding actions over $t-1$ periods. The initial history is empty (trivial) and denoted by $H^1 = \emptyset$. Also, let $H^\infty = \cup_{t=1}^\infty H^t$. A typical history of mechanisms and message profiles played is denoted by $h \in H^\infty$.

A regime, R , is then a mapping, or a set of *transition rules*, $R : H^\infty \rightarrow G$. Let $R|h$ refer to the *continuation regime* that regime R induces at history $h \in H^\infty$. Thus,

$R|h(h') = R(h, h')$ for any $h, h' \in H^\infty$. A regime R is *history-independent* if and only if, for any t and any $h, h' \in H^t$, $R(h) = R(h') \in G$. Notice that, in such a history-independent regime, the specified mechanisms may change over time in a pre-determined sequence. We say that a regime R is *stationary* if and only if, for any $h, h' \in H^\infty$, $R(h) = R(h') \in G$.

Given a regime, a (pure) strategy for an agent depends on the sequences of realized states as well as the histories of mechanisms and message profiles played.⁹ Define $\mathbf{H}^t$ as the $(t - 1)$ -fold Cartesian product of the set $\mathcal{E} \times \Theta$, and let $\mathbf{H}^1 = \emptyset$ and $\mathbf{H}^\infty = \cup_{t=1}^\infty \mathbf{H}^t$ with its typical element denoted by $\mathbf{h}$. Then, each agent i 's corresponding strategy, σ_i , is a mapping $\sigma_i : \mathbf{H}^\infty \times G \times \Theta \rightarrow \cup_{g \in G} M_i^g$ such that $\sigma_i(\mathbf{h}, g, \theta) \in M_i^g$ for any $(\mathbf{h}, g, \theta) \in \mathbf{H}^\infty \times G \times \Theta$. Let Σ_i be the set of all such strategies, and let $\Sigma \equiv \Sigma_1 \times \dots \times \Sigma_I$. A strategy profile is denoted by $\sigma \in \Sigma$. We say that σ_i is a Markov strategy if and only if $\sigma_i(\mathbf{h}, g, \theta) = \sigma_i(\mathbf{h}', g, \theta)$ for any $\mathbf{h}, \mathbf{h}' \in \mathbf{H}^\infty$, $g \in G$ and $\theta \in \Theta$. A strategy profile $\sigma = (\sigma_1, \dots, \sigma_I)$ is Markov if and only if σ_i is Markov for each i .

Next, let $\theta(t) = (\theta^1, \dots, \theta^{t-1}) \in \Theta^{t-1}$ denote a sequence of realized states up to, but not including, period t with $\theta(1) = \emptyset$. Let $q(\theta(t)) \equiv p(\theta^1) \times \dots \times p(\theta^{t-1})$. Suppose that R is the regime and σ the strategy profile chosen by the agents. Let us define the following variables *on the outcome path*:

- $\mathbf{h}(\theta(t), \sigma, R) \in \mathbf{H}^t$ denotes the $t - 1$ period history generated by σ in R over state realizations $\theta(t) \in \Theta^{t-1}$.
- $g^{\theta(t)}(\sigma, R) \equiv (M^{\theta(t)}(\sigma, R), \psi^{\theta(t)}(\sigma, R))$ refers to the mechanism played at $\mathbf{h}(\theta(t), \sigma, R)$.
- $m^{\theta(t), \theta^t}(\sigma, R) \in M^{\theta(t)}(\sigma, R)$ refers to the message profile reported at $\mathbf{h}(\theta(t), \sigma, R)$ when the current state is θ^t .
- $a^{\theta(t), \theta^t}(\sigma, R) \equiv \psi^{\theta(t)}(m^{\theta(t), \theta^t}(\sigma, R)) \in A$ refers to the outcome implemented at $\mathbf{h}(\theta(t), \sigma, R)$ when the current state is θ^t .
- $\pi_i^{\theta(t)}(\sigma, R)$, with slight abuse of notation, denotes agent i 's continuation payoff at $\mathbf{h}(\theta(t), \sigma, R)$; that is,

$$\pi_i^{\theta(t)}(\sigma, R) = (1 - \delta) \sum_{s \in \mathbb{Z}_{++}} \sum_{\theta(s) \in \Theta^{s-1}} \sum_{\theta^s \in \Theta} \delta^{s-1} q(\theta(s), \theta^s) u_i(a^{\theta(t), \theta(s), \theta^s}(\sigma, R), \theta^s).$$

⁹Although we restrict our attention to pure strategies, it is possible to extend the analysis to allow for mixed strategies. See Section 5 below.

For notational simplicity, let $\pi_i(\sigma, R) \equiv \pi_i^{\theta(1)}(\sigma, R)$. Also, when the meaning is clear, we shall sometimes suppress the arguments in the above variables and refer to them simply as $\mathbf{h}(\theta(t))$, $g^{\theta(t)}$, $m^{\theta(t), \theta^t}$, $a^{\theta(t), \theta^t}$ and $\pi_i^{\theta(t)}$.

A strategy profile $\sigma = (\sigma_1, \dots, \sigma_I)$ is a Nash equilibrium of regime R if, for each i , $\pi_i(\sigma, R) \geq \pi_i(\sigma'_i, \sigma_{-i}, R)$ for all $\sigma'_i \in \Sigma_i$. Let $\Omega^\delta(R) \subseteq \Sigma$ denote the set of (pure strategy) Nash equilibria of regime R with discount factor δ .

We are now ready to define the following notions of Nash repeated implementation.

Definition 3 *An SCF f is payoff-repeated-implementable in Nash equilibrium from period τ if there exists a regime R such that (i) $\Omega^\delta(R)$ is non-empty; and (ii) every $\sigma \in \Omega^\delta(R)$ is such that $\pi_i^{\theta(t)}(\sigma, R) = v_i(f)$ for any i , $t \geq \tau$ and $\theta(t)$. An SCF f is repeated-implementable in Nash equilibrium from period τ if, in addition, every $\sigma \in \Omega^\delta(R)$ is such that $a^{\theta(t), \theta^t}(\sigma, R) = f(\theta^t)$ for any $t \geq \tau$, $\theta(t)$ and θ^t .*

The first notion represents repeated implementation in terms of payoffs, while the second asks for repeated implementation of outcomes and, therefore, is a stronger concept. Repeated implementation from some period τ requires the existence of a regime in which every Nash equilibrium delivers the correct continuation payoff profile or the correct outcomes from period τ onwards for every possible sequence of state realizations.

3.3 Main results

3.3.1 Necessity

As illustrated by the motivating example in Section 3.1, our understanding of repeated games suggests that some form of efficiency ought to play a necessary role towards repeated implementation. Our first result formalizes this by showing that, if the agents are sufficiently patient and an SCF f is repeated-implementable from any period, then there cannot be another SCF whose *range* also belongs to that of f such that all agents strictly prefer it to f in expectation. Otherwise, there must be a “collusive” equilibrium in which the agents obtain higher payoffs; but this is a contradiction.

Theorem 1 *Consider any SCF f such that $v(f)$ is strictly Pareto dominated by another payoff profile $v' \in V(f)$. Then there exists $\bar{\delta} \in (0, 1)$ such that, for any $\delta \in (\bar{\delta}, 1)$ and*

period τ , f cannot be repeated-implementable in Nash equilibrium from period τ .¹⁰

Proof. Let $\bar{\delta} = \frac{2\rho}{2\rho + \max_i [v'_i - v_i(f)]}$, where $\rho \equiv \max_{i \in I, \theta \in \Theta, a, a' \in A} [u_i(a, \theta) - u_i(a', \theta)]$. Fix any $\delta \in (\bar{\delta}, 1)$. We prove the claim by contradiction. So, suppose that there exists a regime R^* that repeated-implements f from some period τ .

For any strategy profile σ , any player i , any date t and sequence of states $\theta(t)$ and θ^t , let $M_i(\theta(t), \sigma, R^*)$ denote the set of messages that i can play at history $\mathbf{h}(\theta(t), \sigma, R^*)$. Also, with some abuse of notation, for any $m_i \in M_i(\theta(t), \sigma, R^*)$, let $\pi_i^{\theta(t), \theta^t}(\sigma)|m_i$ represent i 's continuation payoff from period $t+1$ if the sequence of states $(\theta(t), \theta^t)$ is observed, i deviates from σ_i for only one period at $\mathbf{h}(\theta(t), \sigma, R^*)$ after observing θ^t and every other agent plays the regime according to σ_{-i} .

Consider any $\sigma^* \in \Omega^\delta(R^*)$. Since σ^* is a Nash equilibrium that repeated-implements f from period τ , the following must be true about the equilibrium path: for any i , $t \geq \tau$, $\theta(t)$, θ^t and $m'_i \in M_i(\theta(t), \sigma^*, R^*)$,

$$(1 - \delta)u_i(a^{\theta(t), \theta^t}(\sigma^*, R^*), \theta^t) + \delta v_i(f) \geq (1 - \delta)u_i(a, \theta^t) + \delta \pi_i^{\theta(t), \theta^t}(\sigma^*)|m'_i,$$

where $a \equiv \psi^{\theta(t)}(m'_i, m_{-i}^{\theta(t), \theta^t}(\sigma^*, R^*))$. This implies that, for any i , $t \geq \tau$, $\theta(t)$, θ^t and $m'_i \in M_i(\theta(t), \sigma^*, R^*)$,

$$\delta \pi_i^{\theta(t), \theta^t}(\sigma^*)|m'_i \leq (1 - \delta)\rho + \delta v_i(f). \quad (1)$$

Next, let $f' \in F(f)$ be the SCF that induces the payoff profile v' . Then, for all i , $v'_i = v_i(f') > v_i(f)$. Also, since $f' \in F(f)$, there must exist a mapping $\lambda : \Theta \rightarrow \Theta$ such that $f'(\theta) = f(\lambda(\theta))$ for all θ . Consider the following strategy profile σ' : for any i , g , and θ , (i) $\sigma'_i(\mathbf{h}, g, \theta) = \sigma_i^*(\mathbf{h}, g, \theta)$ for any $\mathbf{h} \in \mathbf{H}^t$, $t < \tau$; (ii) for any $\mathbf{h} \in \mathbf{H}^t$, $t \geq \tau$, $\sigma'_i(\mathbf{h}, g, \theta) = \sigma_i^*(\mathbf{h}, g, \lambda(\theta))$ if $\mathbf{h}$ is such that there has been no deviation from σ' , while $\sigma'_i(\mathbf{h}, g, \theta) = \sigma_i^*(\mathbf{h}, g, \theta)$ otherwise.

Given the construction of σ' , and since $\sigma^* \in \Omega^\delta(R^*)$ and $\pi_i^{\theta(t)}(\sigma', R) = v'_i > v_i$ for all i , $t \geq \tau$ and $\theta(t)$, no agent wants to deviate from σ' at any history before period τ . Next, fix any i , $t \geq \tau$, $\theta(t)$ and θ^t . By the construction of σ' , and since σ^* repeated-implements f from τ , we also have that agent i 's continuation payoff from σ' at $\mathbf{h}(\theta(t), \sigma', R^*)$ after observing θ^t is given by

$$(1 - \delta)u_i(a^{\theta(t), \theta^t}(\sigma', R^*), \theta^t) + \delta v_i(f'). \quad (2)$$

¹⁰Note that the necessary condition implied by this statement would correspond to *efficiency in the range* if $V(f)$ was a convex set (which would be true, for instance, if public randomization were allowed).

On the other hand, the corresponding payoff from any unilateral one-period deviation $m'_i \in M_i(\theta(t), \sigma', R^*)$ by i from σ' is

$$(1 - \delta)u_i\left(\psi^{\theta(t)}(m'_i, m_{-i}^{\theta(t), \theta^t}(\sigma', R^*)), \theta^t\right) + \delta\pi_i^{\theta(t), \theta^t}(\sigma')|m'_i. \quad (3)$$

Notice that, by the construction of σ' , there exists some $\tilde{\theta}(t)$ such that $\mathbf{h}(\theta(t), \sigma', R^*) = \mathbf{h}(\tilde{\theta}(t), \sigma^*, R^*)$ and, hence, $M_i(\theta(t), \sigma', R^*) = M_i(\tilde{\theta}(t), \sigma^*, R^*)$. Moreover, after a deviation, σ' induces the same continuation strategies as σ^* . Thus, we have

$$\pi_i^{\theta(t), \theta^t}(\sigma')|m'_i = \pi_i^{\tilde{\theta}(t), \lambda(\theta^t)}(\sigma^*)|m'_i.$$

Then, by (1) above, the deviation payoff (3) is less than or equal to

$$(1 - \delta) \left[u_i\left(\psi^{\theta(t)}(m'_i, m_{-i}^{\theta(t), \theta^t}(\sigma', R^*)), \theta^t\right) + \rho \right] + \delta v_i(f).$$

This, together with $v_i(f') > v_i(f)$, $\delta > \bar{\delta}$ and the definition of $\bar{\delta}$, implies that (2) exceeds (3). But, this means that $\sigma' \in \Omega^\delta(R^*)$. Since σ' induces an average payoff $v_i(f') \neq v_i(f)$ for all i from period τ , we then have a contradiction against the assumption that R^* repeatedly-implements f from τ . ■

3.3.2 Sufficiency

Let us now investigate if an efficient SCF can indeed be repeatedly implemented. We begin with some additional definitions and an important general observation.

First, we call a *trivial* mechanism one that enforces a single outcome. Formally, $\phi(a) = (M, \psi)$ is such that $M_i = \{\emptyset\}$ for all i and $\psi(m) = a \in A$ for all $m \in M$. Also, let $d(i)$ denote a *dictatorial* mechanism in which agent i is the dictator; formally, $d(i) = (M, \psi)$ is such that $M_i = A$, $M_j = \{\emptyset\}$ for all $j \neq i$ and $\psi(m) = m_i$ for all $m \in M$.

Next, let $v_i^i = \sum_{\theta \in \Theta} p(\theta) \max_{a \in A} u_i(a, \theta)$ denote agent i 's maximal one-period payoff. Clearly, v_i^i is i 's payoff when i is the dictator and he acts rationally. Also, let $A^i(\theta) \equiv \{\arg \max_{a \in A} u_i(a, \theta)\}$ represent the set of i 's best outcomes in state θ . Then, define the maximum payoff i can obtain when agent $j \neq i$ is the dictator by $v_i^j = \sum_{\theta \in \Theta} p(\theta) \max_{a \in A^j(\theta)} u_i(a, \theta)$.

We make the following assumption throughout the paper.

- (A) There exist some i and j such that $A^i(\theta) \cap A^j(\theta)$ is empty for some θ .

This assumption is equivalent to assuming that $v_i^i \neq v_i^j$ for some i and j . It implies that in some state there is a conflict between some agents on the best outcome. Since we are concerned with repeated implementation of efficient SCFs, Assumption (A) incurs no loss of generality when each agent has a unique best outcome for each state: if Assumption (A) were not to hold, we could then simply let any agent choose the outcome in each period to obtain repeated implementation of an efficient SCF.

Our results on efficient repeated implementation below are based on the following relatively innocuous auxiliary condition.

Condition ω . For each i , there exists some $\tilde{a}^i \in A$ such that $v_i(f) \geq v_i(\tilde{a}^i)$.

This property says that, for *each* agent, the expected utility that he derives from the SCF is bounded below by that of some constant SCF.¹¹ Note that the property does not require that there be a single constant SCF to provide the lower bound for all agents. In many applications, condition ω is naturally satisfied.

Now, let Φ^a denote a stationary regime in which the trivial mechanism $\phi(a)$ is repeated forever and let D^i denote a stationary regime in which the dictatorial mechanism $d(i)$ is repeated forever. Also, let $\mathcal{S}(i, a)$ be the set of all possible history-independent regimes in which the enforced mechanisms are either $d(i)$ or $\phi(a)$ only. For any $i, j \in I$, $a \in A$ and $S^i \in \mathcal{S}(i, a)$, we denote by $\pi_j(S^i)$ the maximum payoff j can obtain when S^i is enforced and agent i always chooses a best outcome in the dictatorial mechanism $d(i)$.¹²

Our first Lemma applies the result of Sorin [27] to our setup. If an SCF satisfies condition ω , any individual's corresponding payoff can be generated precisely by a sequence of appropriate dictatorial and trivial mechanisms, as long as the discount factor is greater than a half.

Lemma 1 *Consider an SCF f and any i . Suppose that there exists some $\tilde{a}^i \in A$ such that $v_i(f) \geq v_i(\tilde{a}^i)$. Then, for any $\delta \geq \frac{1}{2}$, there exists a regime $S^i \in \mathcal{S}(i, \tilde{a}^i)$ such that $\pi_i(S^i) = v_i(f)$.*

Proof. By assumption there exists some outcome $\tilde{a}^i$ such that $v_i(f) \in [v_i(\tilde{a}^i), v_i^i]$. Since $v_i(\tilde{a}^i)$ is the one-period payoff of i when $\phi(a)$ is played and v_i^i is i 's payoff when $d(i)$ is

¹¹We later discuss an alternative requirement, which we call *non-exclusion*, that can serve the same purpose as condition ω in our analysis.

¹²Note when the best choice of i is not unique in some state then the payoff of any $j \neq i$ may depend on the precise choice of i when i is the dictator.

played and i behaves rationally, it follows from the algorithm of Sorin [27] (see Lemma 3.7.1 of Mailath and Samuelson [16]) that there exists a regime $S^i \in \mathcal{S}(i, \tilde{a}^i)$ that alternates between $\phi(a)$ and $d(i)$, and generates the payoff $v_i(f)$ exactly. ■

The above statement assumes that $\delta \geq \frac{1}{2}$ because $v_i(f)$ is a convex combination of exactly *two* payoffs $v_i(\tilde{a}^i)$ and v_i^i . For the remainder of the paper, unless otherwise stated, δ will be fixed to be no less than $\frac{1}{2}$ as required by this Lemma. But, note that if the environment is sufficiently rich that, for each i , one can find some $\tilde{a}^i$ with $v_i(\tilde{a}^i) = v_i(f)$ (for instance, when utilities are transferable) then our results below are true for any $\delta \in (0, 1)$.

Three or more agents The analysis with three or more agents is somewhat different from that with two players. Here, we consider the former case and assume that $I > 2$.

Our arguments are constructive. First, fix any SCF f that satisfies condition ω and define mechanism $g^* = (M, \psi)$ as follows: $M_i = \Theta \times \mathbb{Z}_+$ for all i , and ψ is such that (i) if $m_i = (\theta, \cdot)$ for at least $I - 1$ agents, $\psi(m) = f(\theta)$ and (ii) if $m = ((\theta^i, z^i))_{i \in I}$ is of any other type, $\psi(m) = f(\tilde{\theta})$ for some arbitrary but fixed state $\tilde{\theta} \in \Theta$.

Next, let R^* denote any regime satisfying the following transition rules: $R^*(\emptyset) = g^*$ and, for any $h = ((g^1, m^1), \dots, (g^{t-1}, m^{t-1})) \in H^t$ such that $t > 1$ and $g^{t-1} = g^*$,

1. if $m_i^{t-1} = (\cdot, 0)$ for all i , $R^*(h) = g^*$;
2. if there exists some i such that $m_j^{t-1} = (\cdot, 0)$ for all $j \neq i$ and $m_i^{t-1} = (\cdot, z^i)$ with $z^i > 0$, $R^*|h = S^i$, where $S^i \in \mathcal{S}(i, \tilde{a}^i)$ such that $v_i(\tilde{a}^i) \leq v_i(f)$ and $\pi_i(S^i) = v_i(f)$ (by condition ω and Lemma 1, regime S^i exists);
3. if m^{t-1} is of any other type and i is lowest-indexed agent among those who announce the highest integer, $R^*|h = D^i$.

Regime R^* starts with mechanism g^* . At any period in which this mechanism is played, the transition is as follows. If all agents announce zero integer, then the mechanism next period continues to be g^* . If all agents but one, say i , announce zero integer and i does not, then the continuation regime at the next period is a history-independent regime in which the odd-one-out i can guarantee himself a payoff exactly equal to the target level $v_i(f)$ (invoking Lemma 1). Finally, if the message profile is of any other type, one of the agents who announce the highest integer becomes a dictator forever thereafter.

Note that, unless all agents “agree” on zero integer when playing mechanism g^* , any strategic play in regime R^* effectively ends; for any other message profile, the continuation regime is history-independent and employs only dictatorial and/or trivial mechanisms.

We now characterize the set of Nash equilibria of regime R^* . A critical feature of our regime construction is conveyed in our next Lemma: beyond the first period, as long as g^* is the mechanism played, each agent i ’s equilibrium continuation payoff is always bounded below by the target payoff $v_i(f)$. This follows from $\pi_i(S^i) = v_i(f)$ (Lemma 1).

Lemma 2 *Suppose that f satisfies condition ω . Fix any $\sigma \in \Omega^\delta(R^*)$. For any $t > 1$ and $\theta(t)$, if $g^{\theta(t)}(\sigma, R^*) = g^*$, $\pi_i^{\theta(t)}(\sigma, R^*) \geq v_i(f)$ for all i .*

Proof. Suppose not; then, at some $t > 1$ and $\theta(t)$, $\pi_i^{\theta(t)}(\sigma, R^*) < v_i(f)$ for some i . Let $\theta(t) = (\theta(t-1), \theta^{t-1})$. By the transition rules of R^* , it must be that $g^{\theta(t-1)}(\sigma, R^*) = g^*$ and, for all i , $m_i^{\theta(t-1), \theta^{t-1}}(\sigma, R^*) = (\theta_i, 0)$ for some θ_i .

Consider agent i deviating to another strategy σ'_i identical to the equilibrium strategy σ_i at every history, except at $\mathbf{h}(\theta(t-1), \sigma, R^*)$ and period $t-1$ state θ^{t-1} where it announces the state announced by σ_i , θ_i , and a positive integer. Given ψ of mechanism g^* , the outcome at $(\mathbf{h}(\theta(t-1), \sigma, R^*), \theta^{t-1})$ does not change, $a^{\theta(t-1), \theta^{t-1}}(\sigma'_i, \sigma_{-i}, R^*) = a^{\theta(t-1), \theta^{t-1}}(\sigma, R^*)$, while, by transition rule 2 of R^* , regime S^i will be played thereafter and i can obtain continuation payoff $v_i(f)$ at the next period. Thus, the deviation is profitable, contradicting the Nash equilibrium assumption. ■

We next want to show that indeed mechanism g^* will always be played on the equilibrium path. To this end, we also assume that $\tilde{a}^i \in A$ used in the construction of S^i (in the definition of R^* above) is such that

$$\text{if } v_i(f) = v_i(\tilde{a}^i) \text{ then } v_j^j > v_j(\tilde{a}^i) \text{ for some } j. \quad (4)$$

Condition (4) ensures that, in any equilibrium, agents will always agree and, therefore, g^* will always be played, implementing outcomes belonging only to the *range* of f .

Lemma 3 *Suppose that f satisfies ω . Also, suppose that, for each i , outcome $\tilde{a}^i \in A$ used in the construction of S^i above satisfies (4). Then, for any $\sigma \in \Omega^\delta(R^*)$, t , $\theta(t)$ and θ^t , we have: (i) $g^{\theta(t)}(\sigma, R^*) = g^*$; (ii) $m_i^{\theta(t), \theta^t}(\sigma, R^*) = (\cdot, 0)$ for all i ; (iii) $a^{\theta(t), \theta^t}(\sigma, R^*) \in f(\Theta)$.*

Proof. First we establish two claims.

Claim 1: Fix any i and any $a^i(\theta) \in A^i(\theta)$ for every θ . There exists $j \neq i$ such that $v_j^j > \sum_{\theta} p(\theta) u_j(a^i(\theta), \theta)$.

To prove this claim, suppose otherwise; then $v_j^j = \sum_{\theta} p(\theta) u_j(a^i(\theta), \theta)$ for all $j \neq i$. But this means that $a^i(\theta) \in A^j(\theta)$ for all $j \neq i$ and θ . Since by assumption $a^i(\theta) \in A^i(\theta)$, this contradicts Assumption (A).

Claim 2: Fix any $\sigma \in \Omega^{\delta}(R^)$, t , $\theta(t)$ and θ^t . If $g^{\theta(t)} = b^*$ and $m_i^{\theta(t), \theta^t} = (\cdot, z^i)$ with $z^i > 0$ for some i then there must exist some $j \neq i$ such that $\pi_j^{\theta(t), \theta^t} < v_j^j$.*

To prove this claim note that, given the definition of R^* , the continuation regime at the next period is either D^i or S^i for some i . There are two cases to consider.

Case 1: The continuation regime is $S^i = \Phi^{\tilde{a}^i}$ (S^i enforces $\tilde{a}^i \in A$ at every period).

In this case $\pi_i^{\theta(t), \theta^t} = v_i(f) = v_i(\tilde{a}^i)$. Then the claim follows from $\pi_j^{\theta(t), \theta^t} = v_j(\tilde{a}^i)$ and condition (4).

Case 2: The continuation regime is either D^i or $S^i \neq \Phi^{\tilde{a}^i}$.

By assumption under $d(i)$ every agent j receives at most $v_j^i \leq v_j^j$. Also, when the trivial mechanism $\phi(\tilde{a}^i)$ is played every agent j receives a payoff $v_j(\tilde{a}^i) \leq v_j^j$. Since in this case the continuation regime involves playing either $d(i)$ or $\phi(\tilde{a}^i)$, it follows that, for every j , $\pi_j^{\theta(t), \theta^t} \leq v_j^j$. Furthermore, by Claim 1, it must be that this inequality is strict for some $j \neq i$. This is because in this case there exists some $t' > t$ and some sequence of states $\theta(t') = (\theta(t), \theta^{t+1}, \dots, \theta^{t'-1})$ such that the continuation regime enforces $d(i)$ at history $\mathbf{h}(\theta(t'))$; but then $a^{\theta(t'), \theta} \in A^i(\theta)$ for all θ and therefore, by Claim 1, there exists an agent $j \neq i$ such that $v_j^j > \sum_{\theta} p(\theta) u_j(a^{\theta(t'), \theta}, \theta)$.

Next note that, given the definitions of g^* and R^* , to prove the lemma it suffices to show the following: *For any t and $\theta(t)$, if $g^{\theta(t)} = g^*$ then $m_i^{\theta(t), \theta^t} = (\cdot, 0)$ for all i and θ^t .*

To show this, suppose otherwise; so, at some t and $\theta(t)$, $g^{\theta(t)} = b^*$ but $m_i^{\theta(t), \theta^t} = (\cdot, z^i)$ with $z^i > 0$ for some i and θ^t . Then by Claim 2 there exists $j \neq i$ such that $\pi_j^{\theta(t), \theta^t} < v_j^j$. Next consider j deviating to another strategy which yields the same outcome path as the equilibrium strategy, σ_j , at every history, except at $(\mathbf{h}(\theta(t)), \theta^t)$ where it announces the state announced by σ_i and an integer higher than any integer that can be reported by σ at this history. Given ψ , such a deviation does not incur a one-period utility loss while strictly improving the continuation payoff as of the next period since the deviator j becomes a dictator himself and by the previous argument $\pi_j^{\theta(t), \theta^t} < v_j^j$. This is a contradiction. ■

Given the previous two lemmas, we can now pin down the equilibrium payoffs by invoking efficiency *in the range*.

Lemma 4 Suppose that f is efficient in the range and satisfies condition ω . Also, suppose that, for each i , outcome $\tilde{a}^i \in A$ used in the construction of S^i above satisfies (4). Then, for any $\sigma \in \Omega^\delta(R^*)$, $\pi_i^{\theta(t)}(\sigma, R^*) = v_i(f)$ for any i , $t > 1$ and $\theta(t)$.

Proof. Suppose not; then f is efficient in the range but there exist some $\sigma \in \Omega^\delta(R^*)$, $t > 1$ and $\theta(t)$ such that $\pi_i^{\theta(t)} \neq v_i(f)$ for some i . By Lemma 2, it must be that $\pi_i^{\theta(t)} > v_i(f)$. Also, by part (iii) of Lemma 3, $\left(\pi_j^{\theta(t)}\right)_{j \in I} \in co(V(f))$. Since f is efficient in the range, it then follows that there must exist some $j \neq i$ such that $\pi_j^{\theta(t)} < v_j(f)$. But, this contradicts Lemma 2. ■

It is straightforward to establish that regime R^* has a Nash equilibrium in Markov strategies which attains truth-telling and, hence, the desired social choice at every possible history.

Lemma 5 Suppose that f satisfies condition ω . There exists $\sigma^* \in \Omega^\delta(R^*)$, which is Markov, such that, for any t , $\theta(t)$ and θ^t , (i) $g^{\theta(t)}(\sigma^*, R^*) = g^*$; (ii) $a^{\theta(t), \theta^t}(\sigma^*, R^*) = f(\theta^t)$.

Proof. Consider $\sigma^* \in \Sigma$ such that, for all i , $\sigma_i^*(\mathbf{h}, g^*, \theta) = \sigma_i^*(\mathbf{h}', g^*, \theta) = (\theta, 0)$ for any $\mathbf{h}, \mathbf{h}' \in \mathbf{H}^\infty$ and θ . Thus, at any t and $\theta(t)$, we have $\pi_i^{\theta(t)}(\sigma^*, R^*) = v_i(f)$ for all i . Consider any i making a unilateral deviation from σ^* by choosing some $\sigma'_i \neq \sigma_i^*$ which announces a different message at some $(\theta(t), \theta^t)$. But, given ψ of g^* , it follows that $a^{\theta(t), \theta^t}(\sigma'_i, \sigma_{-i}^*, R^*) = a^{\theta(t), \theta^t}(\sigma^*, R^*) = f(\theta^t)$ while, by transition rule 2 of R^* , $\pi_i^{\theta(t), \theta^t}(\sigma'_i, \sigma_{-i}^*, R^*) = v_i(f)$. Thus, the deviation is not profitable.¹³ ■

We are now ready to present our main results for the complete information case. The first result requires a slight strengthening of condition ω in order to ensure implementation of SCFs that are efficient *in the range*.

Condition ω' . For each i , there exists some $\tilde{a}^i \in A$ such that

- (a) $v_i(f) \geq v_i(\tilde{a}^i)$ and
- (b) if $v_i(f) = v_i(\tilde{a}^i)$ then either (i) there exists j such that $v_j^j > v_j(\tilde{a}^i)$ or (ii) the payoff profile $v(\tilde{a}^i)$ does not Pareto dominate $v(f)$.

¹³In the Nash equilibrium constructed for R^* , each agent is indifferent between the equilibrium and any unilateral deviation. The following modification to regime R^* will admit a *strict* Nash equilibrium with the same properties: for each i , construct S^i such that i obtains a payoff $v_i(f) - \epsilon$ for some arbitrarily small $\epsilon > 0$. This will, however, result in the equilibrium payoffs of our canonical regime to *approximate* the efficient target payoffs.

Note that the additional requirement (b) in ω' that does not appear in condition ω applies only if $v_i(f) = v_i(\tilde{a}^i)$.

Theorem 2 *Suppose that $I \geq 3$, and consider an SCF f satisfying condition ω' . If f is efficient in the range, it is payoff-repeated-implementable in Nash equilibrium from period 2; if f is strictly efficient in the range, it is repeated-implementable in Nash equilibrium from period 2.*

Proof. Consider any profile of outcomes $(\tilde{a}^1, \dots, \tilde{a}^I)$ satisfying condition ω' . There are two cases to consider.

Case 1: For all i , $\tilde{a}^i \in A$ satisfies (4). In this case the first part of the theorem follows immediately from Lemmas 4 and 5.

To prove the second part, fix any $\sigma \in \Omega^\delta(R^*)$, i , $t > 1$ and $\theta(t)$. Then

$$\pi_i^{\theta(t)} = \sum_{\theta^t \in \Theta} p(\theta^t) \left[(1 - \delta) u_i(a^{\theta(t), \theta^t}, \theta^t) + \delta \pi_i^{\theta(t), \theta^t} \right]. \quad (5)$$

Also, by Lemma 4 we have $\pi_i^{\theta(t)} = v_i(f)$ and $\pi_i^{\theta(t), \theta^t} = v_i(f)$ for any θ^t . But then, by (5), we have $\sum_{\theta^t} p(\theta^t) u_i(a^{\theta(t), \theta^t}, \theta^t) = v_i(f)$. Since, by part (iii) of Lemma 3, $a^{\theta(t), \theta^t} \in f(\Theta)$, and since f is strictly efficient in the range, the claim follows.

Case 2: For some i , condition (4) does not hold.

In this case, $v_i(f) = v_i(\tilde{a}^i)$ and $v_j^j = v_j(\tilde{a}^i)$ for all $j \neq i$. Then, by b(ii) of condition ω' , $v(\tilde{a}^i)$ does not Pareto dominate $v(f)$. Since $v_j^j \geq v_j(f)$, it must then be that $v_j(f) = v_j(\tilde{a}^i)$ for all j . Such an SCF can be trivially payoff-repeated-implemented via $\Phi^{\tilde{a}^i}$. Furthermore, since $v_j(f) = v_j(\tilde{a}^i) = v_j^j$ for all j , f is efficient. Thus, if f is strictly efficient (in the range), it must be constant, i.e. $f(\theta) = \tilde{a}^i$ for all θ , and hence can also be repeated-implemented via $\Phi^{\tilde{a}^i}$. ■

Note that when f is efficient (over the entire set of SCFs) part b(ii) of condition ω' is vacuously satisfied. Therefore, we can use condition ω instead of ω' to establish repeated implementation with efficiency.

Corollary 1 *Suppose that $I \geq 3$, and consider an SCF f satisfying condition ω . If f is efficient, it is payoff-repeated-implementable in Nash equilibrium from period 2; if f is strictly efficient, it is repeated-implementable in Nash equilibrium from period 2.*

Note that Theorem 2 and its Corollary establish repeated implementation from the second period and, therefore, unwanted outcomes may still be implemented in the first period. This point will be discussed in more detail in Section 5 below.

Two agents As in one-shot Nash implementation (Moore and Repullo [24] and Dutta and Sen [10]), the two-agent case brings non-trivial differences to the analysis. In particular, with three or more agents a unilateral deviation from “consensus” can be detected; with two agents, however, it is not possible to identify the misreport in the event of disagreement. In our repeated implementation setup, this creates a difficulty in establishing existence of an equilibrium in the canonical regime.

As identified by Dutta and Sen [10], a necessary condition for existence of an equilibrium in the one-shot setup is a *self-selection* requirement that ensures the availability of a punishment whenever the two players disagree on their announcements of the state but one of them is telling the truth. We show below that, with two agents, such a condition together with condition ω , or ω' , delivers repeated implementation of an SCF that is efficient, or efficient in the range. Formally, for any f , i and θ , let $L_i(\theta) = \{a \in A | u_i(a, \theta) \leq u_i(f(\theta), \theta)\}$ be the set of outcomes that are no better than f for agent i . We say that f satisfies *self-selection* if $L_1(\theta) \cap L_2(\theta') \neq \emptyset$ for any $\theta, \theta' \in \Theta$.

Theorem 3 *Suppose that $I = 2$, and consider an SCF f satisfying condition ω (ω') and self-selection. If f is efficient (in the range), it is payoff-repeated-implementable in Nash equilibrium from period 2; if f is strictly efficient (in the range), it is repeated-implementable in Nash equilibrium from period 2.*

The proof of the above theorem is relegated to the Appendix. We note that self-selection and condition ω are weaker than the condition appearing in Moore and Repullo [24] which requires the existence of an outcome that is *strictly* worse for both players in *every* state; with self-selection the requirement is that for each pair of states there exists an outcome such that, in those two states, neither player is better off compared with what each would obtain with f .

A similar result can be obtained with an alternative condition to self-selection. We show in the Supplementary Material that, *with sufficiently patient agents*, the two requirements of self-selection and condition ω' needed to establish repeated implementation of an SCF that is efficient in the range in Theorem 3 above can be replaced by an assumption

that stipulates the existence of an outcome $\tilde{a}$ that is strictly worse than f *on average* for both players: $v_i(\tilde{a}) < v_i(f)$ for all $i = 1, 2$.

4 Incomplete information

4.1 The Setup

We now extend our analysis to incorporate incomplete information. An implementation problem with incomplete information is denoted by $\tilde{\mathcal{P}} = [I, A, \Theta, p, (u_i)_{i \in I}]$, and this modifies the previous definition, $\mathcal{P}$, with the following additions: $\Theta = \prod_{i \in I} \Theta_i$ is a finite set of states, where Θ_i denotes the finite set of agent i 's *types*; let $\theta_{-i} \equiv (\theta_j)_{j \neq i}$ and $\Theta_{-i} \equiv \prod_{j \neq i} \Theta_j$; p denotes a probability distribution defined on Θ (such that $p(\theta) > 0$ for each θ); for each i , let $p_i(\theta_i) = \sum_{\theta_{-i}} p(\theta_{-i}, \theta_i)$ be the marginal probability of type θ_i and $p_i(\theta_{-i}|\theta_i) = p(\theta_{-i}, \theta_i)/p_i(\theta_i)$ be the conditional probability of θ_{-i} given θ_i .

The one-period utility function for i is defined as before by $u_i : \Theta \times A \rightarrow \mathbb{R}$ and the interim expected utility/payoff of outcome a to agent i of type θ_i is given by $v_i(a|\theta_i) = \sum_{\theta_{-i} \in \Theta_{-i}} p_i(\theta_{-i}|\theta_i) u_i(a, \theta_{-i}, \theta_i)$. Ex ante payoff of outcome a is then defined by $v_i(a) = \sum_{\theta_i \in \Theta_i} v_i(a|\theta_i) p(\theta_i)$. Similarly, with slight abuse of notation, define respectively the interim and ex ante payoffs of an SCF $f : \Theta \rightarrow A$ to agent i of type θ_i by $v_i(f|\theta_i) = \sum_{\theta_{-i} \in \Theta_{-i}} p_i(\theta_{-i}|\theta_i) u_i(f(\theta_{-i}, \theta_i), \theta_{-i}, \theta_i)$ and $v_i(f) = \sum_{\theta_i \in \Theta_i} v_i(f|\theta_i) p(\theta_i)$. Also, define the maximum (ex ante) payoff that agent i can obtain if he were a dictator by

$$v_i^i = \sum_{\theta_i \in \Theta_i} p_i(\theta_i) \left[\max_{a \in A} \sum_{\theta_{-i} \in \Theta_{-i}} p_i(\theta_{-i}|\theta_i) u_i(a, \theta_{-i}, \theta_i) \right].$$

Note that as in the complete information case $v_i^i \geq v_i(f)$ for all i and $f \in F$. The definition of an efficient SCF remains as before.

An infinitely repeated implementation problem, $\tilde{\mathcal{P}}^\infty$, represents infinite repetitions of $\tilde{\mathcal{P}} = [I, A, \Theta, p, (u_i)_{i \in I}]$. In each period, the state is drawn from Θ from an independent and identical probability distribution p ; each agent observes only his own type. As in the complete information case, each agent evaluates an infinite outcome sequence, $a^\infty = (a^t, \theta^t)_{t \in \mathbb{Z}_{++}, \theta^t \in \Theta}$, according to discounted average expected utilities.

As before, let $\theta(t) = (\theta^1, \dots, \theta^{t-1}) \in \Theta^{t-1}$ denote a sequence of realized states up to, but not including, period t with $\theta(1) = \emptyset$, and $q(\theta(t)) \equiv p(\theta^1) \times \dots \times p(\theta^{t-1})$. The

definition of a regime also remains as before, with H^∞ denoting the set of all possible publicly observable histories (of mechanisms and messages). To define a strategy, let $\mathbf{H}_i^t = (\mathcal{E} \times \Theta_i)^{t-1}$, where $\mathbf{H}_i^1 = \emptyset$, and $\mathbf{H}_i^\infty = \cup_{t=1}^\infty \mathbf{H}_i^t$ with its typical element denoted by $\mathbf{h}_i$. Then, for any regime R , each agent i 's strategy, σ_i , is a mapping $\sigma_i : \mathbf{H}_i^\infty \times G \times \Theta_i \rightarrow \cup_{g \in G} M_i^g$ such that $\sigma_i(\mathbf{h}_i, g, \theta_i) \in M_i^g$ for any $(\mathbf{h}_i, g, \theta_i) \in \mathbf{H}_i^\infty \times G \times \Theta_i$. Let Σ_i be the set of all such strategies and $\Sigma \equiv \Sigma_1 \times \cdots \times \Sigma_I$. A Markov strategy (profile) can be defined similarly as before. Note that here we are considering the general case with *private* strategies; our results also hold with *public* strategies that depend only on publicly observable histories.

Next, for any regime R and strategy profile $\sigma \in \Sigma$, we define the following on the outcome path:

- $\mathbf{h}_i(\theta(t), \sigma, R) \in \mathbf{H}_i^t$ denotes the $t-1$ period history that agent i observes if all agents play R according to σ and the state/type profile realizations are $\theta(t) \in \Theta^{t-1}$; also let $\mathbf{h}(\theta(t), \sigma, R) = [\mathbf{h}_i(\theta(t), \sigma, R)]_{i \in I}$.
- $g^{\theta(t)}(\sigma, R)$ denotes the mechanism played at $\mathbf{h}(\theta(t), \sigma, R)$.
- $m^{\theta(t), \theta^t}(\sigma, R)$ denotes the message profile reported and $a^{\theta(t), \theta^t}(\sigma, R)$ denotes the outcome implemented at $\mathbf{h}(\theta(t), \sigma, R)$ when the current state is θ^t .
- $\mu_i(\tilde{\theta}(t) | \theta(t), \sigma, R)$ denotes agent i 's posterior *belief* about the other agents' past types, $\tilde{\theta}(t) \in \Theta^{t-1}$, conditional on observing $\mathbf{h}_i(\theta(t), \sigma, R)$ and all playing according to σ . Thus, $\mu_i(\tilde{\theta}(t) | \theta(t), \sigma, R)$ corresponds to the following probability ratio:

$$\Pr(\mathbf{h}_i(\theta(t), \sigma, R) | \tilde{\theta}(t)) q(\tilde{\theta}(t)) / \Pr(\mathbf{h}_i(\theta(t), \sigma, R)).$$

- $E_{\theta(t)} \pi_i^\tau(\sigma, R)$ denotes agent i 's *expected* continuation payoff (i.e. discounted average expected utility) at period $\tau \geq t$ conditional on observing $\mathbf{h}_i(\theta(t), \sigma, R)$ and all playing according to σ . Thus, $E_{\theta(t)} \pi_i^\tau(\sigma, R)$ is given by

$$(1-\delta) \sum_{\tilde{\theta}(t)} \sum_{s \geq \tau} \sum_{\theta(s-t+1)} \sum_{\theta^s} \delta^{s-\tau} q(\theta(s-t+1), \theta^s) \mu_i(\tilde{\theta}(t) | \theta(t), \sigma, R) u_i(a^{\tilde{\theta}(t), \theta(s-t+1), \theta^s}(\sigma, R), \theta^s),$$

where, as before, δ stands for the common discount factor. For simplicity, let $E_{\theta(1)} \pi_i^\tau(\sigma, R) = E \pi_i^\tau(\sigma, R)$ and $E \pi_i^1(\sigma, R) = \pi_i(\sigma, R)$.

When the meaning is clear, we shall sometimes suppress the arguments in the above variables and refer to them simply as $\mathbf{h}_i(\theta(t))$, $g^{\theta(t)}$, $m^{\theta(t), \theta^t}$, $a^{\theta(t), \theta^t}$ and $E_{\theta(t)} \pi_i^\tau$.

4.2 Bayesian repeated implementation

The standard solution concept for implementation with incomplete information is Bayesian Nash equilibrium. In our repeated setup, a strategy profile $\sigma = (\sigma_1, \dots, \sigma_I)$ is a Bayesian Nash equilibrium of R if, for any i , $\pi_i(\sigma, R) \geq \pi_i(\sigma'_i, \sigma_{-i}, R)$ for all σ'_i . Let $Q^\delta(R)$ denote the set of Bayesian Nash equilibria of regime R with discount factor δ . Similarly to the complete information case, we propose the following notions of repeated implementation for the incomplete information case.

Definition 4 *An SCF f is payoff-repeated-implementable in Bayesian Nash equilibrium from period τ if there exists a regime R such that (i) $Q^\delta(R)$ is non-empty; and (ii) every $\sigma \in Q^\delta(R)$ is such that, for every $t \geq \tau$, $E\pi_i^t(\sigma, R) = v_i(f)$ for any i . An SCF f is repeated-implementable in Bayesian Nash equilibrium from period τ if, in addition, every $\sigma \in Q^\delta(R)$ is such that $a^{\theta(t), \theta^t}(\sigma, R) = f(\theta^t)$ for any $\theta(t)$ and θ^t .*

Note that the definition of payoff implementation above is written in terms of expected future payoffs evaluated at the beginning of a regime. This is because we want to avoid the issue of agents' ex post private beliefs that affect their continuation payoffs at different histories.

As in the complete information case, to implement an efficient SCF in Bayesian Nash equilibrium, we assume condition ω : for each i there exists $\tilde{a}^i \in A$ such that $v_i(f) \geq v_i(\tilde{a}^i)$. This condition enables us to extend Lemma 1 above to the incomplete information setup. That is, given any SCF f satisfying condition ω , and any i , one can construct a history-independent regime $S^i \in \mathcal{S}(i, \tilde{a}^i)$ such that $\pi_i(S^i) = v_i(f)$.

In addition to condition ω , here we introduce one additional minor condition.

Condition v. There exist $i, j \in I$, $i \neq j$, such that $v_i(f) < v_i^j$ and $v_j(f) < v_j^i$.

This property requires that there be at least two agents who prefer to be dictators than to have the SCF implemented.

Our main sufficiency result builds on constructing the following regime defined for an SCF f that satisfies condition ω . First, mechanism $b^* = (M, \psi)$ is defined such that (i) for all i , $M_i = \Theta_i \times \mathbb{Z}_+$; and (ii) for any $m = ((\theta_i, z^i))_{i \in I}$, $\psi(m) = f(\theta_1, \dots, \theta_I)$. Then, let B^* represent any regime satisfying the following transition rules: $B^*(\emptyset) = b^*$ and, for any $h = ((g^1, m^1), \dots, (g^{t-1}, m^{t-1})) \in H^t$ such that $t > 1$ and $g^{t-1} = b^*$,

1. if $m_i^{t-1} = (\theta_i, 0)$ for all i , then $B^*(h) = b^*$;
2. if there exists some i such that $m_j^{t-1} = (\theta_j, 0)$ for all $j \neq i$ and $m_i^{t-1} = (\theta_i, z^i)$ with $z^i \neq 0$, then $B^*|h = S^i$, where $S^i \in \mathcal{S}(i, \tilde{a}^i)$ such that $v_i(f) \geq v_i(\tilde{a}^i)$ and $\pi_i(S^i) = v_i(f)$ (by ω and Lemma 1, regime S^i exists);
3. if m^{t-1} is of any other type and i is the lowest-indexed agent among those who announce the highest integer, then $B^*|h = D^i$.

This regime is similar to the regimes constructed for the complete information case. It starts with mechanism b^* in which each agent reports his type and a non-negative integer. Strategic interaction is maintained, and mechanism b^* continues to be played, only if every agent reports zero integer. If an agent reports a non-zero integer, this odd-one-out can obtain a continuation payoff at the next period exactly equal to what he would obtain from implementation of the SCF. If two or more agents report non-zero integers then the one announcing the highest integer becomes a dictator forever as of the next period.

Characterizing the properties of a Bayesian equilibrium of this regime is complicated by the presence of incomplete information, compared to the corresponding task in the complete information setup. First, at any given history, we obtain a lower bound on each agent's *expected* equilibrium continuation payoff *at the next period*. This contrasts with the corresponding result with complete information, Lemma 2, which calculates the lower bound for the actual continuation payoff at any period. With incomplete information, each player i does not know the private information held by others and, therefore, the (off-the-equilibrium) possibility of continuation regime S^i , which guarantees the continuation payoff $v_i(f)$, delivers a lower bound on equilibrium payoff only in expectation.

Lemma 6 *Suppose that f satisfies ω . Fix any $\sigma \in Q^\delta(B^*)$. For any t and $\theta(t)$, if $g^{\theta(t)}(\sigma, B^*) = b^*$, $E_{\theta(t)}\pi_i^{t+1}(\sigma, B^*) \geq v_i(f)$ for all i .*

Proof. Suppose not; so, at some $\theta(t)$, $g^{\theta(t)}(\sigma, B^*) = b^*$ but $E_{\theta(t)}\pi_i^{t+1}(\sigma, B^*) < v_i(f)$ for some i . Then, consider i deviating to another strategy σ'_i identical to the equilibrium strategy σ_i at every history, except at $\mathbf{h}_i(\theta(t), \sigma, B^*)$ where, for each current period realization of θ_i , it reports the same type as σ_i but a different integer which is higher than any integer that can be reported by σ at such a history.

By the definition of b^* , such a deviation does not alter the current period's implemented outcome, regardless of the others' types. As of the next period, it results in either i

becoming a dictator forever (transition rule 3 of B^*) or continuation regime S^i (transition rule 2). Since $v_i^i \geq v_i(f)$ and i can obtain continuation payoff $v_i(f)$ from S^i , the deviation is profitable, implying a contradiction. ■

Second, since we allow for private strategies that depend on each individual's past types, continuation payoffs depend on the agents' posterior beliefs about others' past types and, therefore, a profile of continuation payoffs do not necessarily belong to $co(V)$. Nonetheless, the profile of payoffs evaluated at the beginning of $t = 1$ must belong to $co(V)$ and we can impose efficiency to find upper bounds for such payoffs. From this, it can be shown that continuation payoffs evaluated at any later date must also be bounded above.

Lemma 7 *Suppose that f is efficient and satisfies conditions ω . Fix any $\sigma \in Q^\delta(B^*)$ and any date t . Also, suppose that $g^{\theta(t)}(\sigma, B^*) = b^*$ for all $\theta(t)$. Then, $E_{\theta(t)}\pi_i^{t+1}(\sigma, B^*) = v_i(f)$ for any i and $\theta(t)$.¹⁴*

Proof. Since $g^{\theta(t)} = b^*$ for all $\theta(t)$, it immediately follows from Lemma 6 and Bayes rule that, for all i ,

$$E\pi_i^{t+1} = \sum_{\theta(t)} \Pr(\mathbf{h}_i(\theta(t))) E_{\theta(t)}\pi_i^{t+1} \geq v_i(f). \quad (6)$$

Since $(E\pi_i^{t+1})_{i \in I} \in co(V)$ and f is efficient, (6) then implies that $E\pi_i^{t+1} = v_i(f)$ for all i . By Lemma 6, this in turn implies that $E_{\theta(t)}\pi_i^{t+1} = v_i(f)$ for any i and $\theta(t)$. ■

It remains to be shown that mechanism b^* must always be played along any equilibrium path. Since the regime begins with b^* , we can apply induction to derive this, once next Lemma has been established.

Lemma 8 *Suppose that f is efficient and satisfies conditions ω and v . Fix any $\sigma \in Q^\delta(B^*)$. Also, fix any t , and suppose that $g^{\theta(t)}(\sigma, B^*) = b^*$ for all $\theta(t)$. Then, $g^{\theta(t), \theta^t}(\sigma, B^*) = b^*$ for any $\theta(t)$ and θ^t .*

¹⁴Note that Lemma 7 ties down expected continuation payoffs along the equilibrium path at any period t by assuming that the mechanism played in the corresponding period is b^* for *all* possible state realizations up to t . It imposes no restriction on the mechanisms beyond t . As a result, outcomes that lie outside of the range of f may arise. This is why this Lemma (and thereby all the Bayesian implementation results below) requires the SCF to be efficient and not just efficient in the range.

Proof. Suppose not; so, for some t , $g^{\theta(t)}(\sigma, B^*) = b^*$ for all $\theta(t)$ but $m_i^{\theta(t), \theta^t}(\sigma, B^*) = (\cdot, z)$, $z \neq 0$, for some i , $\theta(t)$ and θ^t . By condition v , there must exist some $j \neq i$ such that $v_j^j > v_j(f)$. Consider j deviating to another strategy identical to the equilibrium strategy, σ_j , except that, at $\mathbf{h}_j(\theta(t), \sigma, B^*)$, it reports the same type as σ_j for each current realization of θ_j but a different integer higher than any integer reported by σ at such a history.

By the definition of b^* , the deviation does not alter the current outcome, regardless of the others' types. But, the continuation regime is D^j if the realized type profile is θ^t while, otherwise, it is D^j or S^j . In the former case, j can obtain continuation payoff $v_j^j > v_j(f)$; in the latter, he can obtain at least $v_j(f)$. Since, by Lemma 7, the equilibrium continuation payoff is $v_j(f)$, the deviation is thus profitable, implying a contradiction. ■

This leads to the following.

Lemma 9 *If f is efficient and satisfies conditions ω and v , every $\sigma \in Q^\delta(B^*)$ is such that $E_{\theta(t)}\pi_i^{t+1}(\sigma, B^*) = v_i(f)$ for any i , t and $\theta(t)$.*

Proof. Since $B^*(\emptyset) = b^*$, it follows by induction from Lemma 8 that, for any t and $\theta(t)$, $g^{\theta(t)}(\sigma, B^*) = b^*$. Lemma 7 then completes the proof. ■

Our objective of Bayesian repeated implementation is now achieved if regime B^* admits an equilibrium. One natural sufficiency condition that will guarantee existence in our setup is incentive compatibility: f is incentive compatible if, for any i and θ_i , $v_i(f|\theta_i) \geq \sum_{\theta_{-i} \in \Theta_{-i}} p_i(\theta_{-i}|\theta_i) u_i(f(\theta_{-i}, \theta_i'), (\theta_{-i}, \theta_i))$ for all $\theta_i' \in \Theta_i$. It is straightforward to see that, if f is incentive compatible, B^* admits a Markov Bayesian Nash equilibrium in which each agent always reports his true type and zero integer. Together with Lemma 9, this immediately leads to our next Theorem.

Theorem 4 *Fix any $I \geq 2$. If f is efficient, incentive compatible and satisfies conditions ω and v , f is payoff-repeated-implementable in Bayesian Nash equilibrium from period 2.*

In the previous section with complete information, to obtain the stronger implementation result in terms of outcomes, we strengthened the efficiency notion by requiring, in addition, that there is no SCF $f' \neq f$ such that $v_i(f) = v_i(f')$ for all i . Here, we need to assume the following to obtain outcome implementation.

Condition χ . There exists no $\gamma : \Theta \times \Theta \rightarrow A$ such that

$$v_i(f) = \sum_{\theta, \theta' \in \Theta} p(\theta)p(\theta')u_i(\gamma(\theta, \theta'), \theta') \text{ for all } i.$$

Corollary 2 *In addition to the conditions in Theorem 4, suppose f satisfies condition χ . Then f is repeated-implementable in Bayesian Nash equilibrium from period 2.*

Proof. Fix any $\sigma \in Q^\delta(B^*)$, i , t and $\theta(t)$. Then we have

$$E_{\theta(t)}\pi_i^{t+1} = \sum_{\theta^t \in \Theta} p(\theta^t) \left[(1 - \delta) \sum_{\theta^{t+1} \in \Theta} p(\theta^{t+1}) u_i(a^{\theta(t), \theta^t, \theta^{t+1}}, \theta^{t+1}) + \delta E_{\theta(t), \theta^t} \pi_i^{t+2} \right]. \quad (7)$$

But, by Lemma 9, $E_{\theta(t)}\pi_i^{t+1} = v_i(f)$ and $E_{\theta(t), \theta^t} \pi_i^{t+2} = v_i(f)$ for any θ^t . Thus, (7) implies that $\sum_{\theta^t, \theta^{t+1}} p(\theta^t)p(\theta^{t+1})u_i(a^{\theta(t), \theta^t, \theta^{t+1}}, \theta^{t+1}) = v_i(f)$. This contradicts condition χ . ■

4.3 More on incentive compatibility

Theorem 4 and its Corollary establish Bayesian repeated implementation of an efficient SCF without (Bayesian) monotonicity but they still assume incentive compatibility to ensure existence of an equilibrium in which every agent always reports his true type. We now explore if it is in fact possible to construct a regime that keeps the desired equilibrium properties and admits such an equilibrium without incentive compatibility.

Interdependent values Many authors have identified a conflict between efficiency and incentive compatibility in the one-shot setup with interdependent values in which some agents' utilities depend on others' private information. (for instance, Maskin [17] and Jehiel and Moldovanu [14]). Thus, it is of particular interest that we establish non-necessity of incentive compatibility for repeated implementation in this case.

Let us assume that the agents know their utilities from the implemented outcomes at the end of each period, and define *identifiability* as follows.

Definition 5 *An SCF f is identifiable if, for any i , $\theta_i, \theta'_i \in \Theta_i$ such that $\theta'_i \neq \theta_i$ and $\theta_{-i} \in \Theta_{-i}$, there exists some $j \neq i$ such that $u_j(f(\theta'_i, \theta_{-i}), \theta'_i, \theta_{-i}) \neq u_j(f(\theta'_i, \theta_{-i}), \theta_i, \theta_{-i})$.*

In words, identifiability requires that, whenever there is one agent lying about his type while all others report their types truthfully, there exists another agent who obtains a (one-period) utility different from what he would have obtained under everyone behaving truthfully.¹⁵ Thus, with an identifiable SCF, if an agent deviates from an equilibrium in which all agents report their types truthfully then there will be at least one other agent who can detect the lie at the end of the period. Notice that the above definition does not require that the detector knows who has lied; he only learns that *someone* has. Identifiability will enable us to build a regime which admits a truth-telling equilibrium based on incentives of repeated play, instead of one-shot incentive compatibility of the SCF. Such incentives involve punishment when someone misreports his type.

To allow the possibility of punishment we need to strengthen condition ω to allow for the existence of an outcome that is strictly worse than the SCF. Specifically, we assume the following.

Condition ω'' . There exists some $\tilde{a} \in A$ such that $v_i(\tilde{a}) < v_i(f)$ for all i .¹⁶

Consider an SCF f that satisfies conditions ω'' and v . Define Z as a mechanism in which (i) for all i , $m_i = \mathbb{Z}_+$; and (ii) for all m , $\psi(m) = a$ for some arbitrary but fixed a , and define $\tilde{b}^*$ as the following extensive form mechanism:

- Stage 1 - Each agent i announces his private information θ_i , and $f(\theta_1, \dots, \theta_I)$ is implemented.
- Stage 2 - Once agents learn their utilities, but before a new state is drawn, each of them announces a report belonging to the set $\{NF, F\} \times \mathbb{Z}_+$, where NF and F refer to “no flag” and “flag” respectively.

The agents’ actions in Stage 2 do not affect the outcome implemented and payoffs in the current period but they determine the continuation play in the regime below.¹⁷ Next define regime $\tilde{B}^*$ to be such that $\tilde{B}^*(\emptyset) = Z$ and satisfies the following transition rules:

1. for any $h = (g^1, m^1) \in H^2$,

¹⁵Notice that identifiability cannot hold with private values.

¹⁶Note that this property subsumes condition ω' .

¹⁷A similar mechanism design is considered by Mezzetti [22] in the one-shot context with interdependent values and quasi-linear utilities.

- (a) if $m_i^1 = (0)$ for all i , then $\tilde{B}^*(h) = \tilde{b}^*$;
 - (b) if there exists some i such that $m_j^1 = (0)$ for all $j \neq i$ and $m_i^1 = z^i$ with $z^i \neq 0$, then $\tilde{B}^*|h = S^i$, where $S^i \in \mathcal{S}(i, \tilde{a}) \setminus \phi^{\tilde{a}}$ such that $v(\tilde{a}) < v(f)$ and $\pi_i(S^i) = v_i(f)$ (by assumption ω'' and Lemma 1 regime S^i exists);
 - (c) if m^1 is of any other type and i is the lowest-indexed agent among those who announce the highest integer, then $\tilde{B}^*|h = D^i$;
2. for any $h = ((g^1, m^1), \dots, (g^{t-1}, m^{t-1})) \in H^t$ such that $t > 2$ and $g^{t-1} = \tilde{b}^*$:
- (a) if m_i^{t-1} is such that every agent reports NF and 0 in Stage 2, then $\tilde{B}^*(h) = \tilde{b}^*$;
 - (b) if m_i^{t-1} is such that at least one agent reports F in Stage 2, then $\tilde{B}^*|h = \Phi^{\tilde{a}}$;
 - (c) if m_i^{t-1} is such that every agent reports NF and every agent except some i announces 0, then $\tilde{B}^*|h = S^i$, where S^i is as in 1(b) above;
 - (d) if m^{t-1} is of any other type and i is the lowest-indexed agent among those who announce the highest integer, then $\tilde{B}^*|h = D^i$.

This regime begins with a simple integer mechanism and non-contingent implementation of an arbitrary outcome in the first period. If all agents report zero, then the next period's mechanism is $\tilde{b}^*$; otherwise, strategic interaction ends with the continuation regime being either S^i or D^i for some i .

The new mechanism $\tilde{b}^*$ sets up two reporting stages. In particular, each agent is endowed with an opportunity to report detection of a lie by raising “flag” (though he may not know who the liar is) after an outcome has been implemented and his own within-period payoff learned. The second stage also features integer play, with the transitions being the same as before as long as every agent reports “no flag”. But, only one “flag” is needed to overrule the integers and activate permanent implementation of outcome $\tilde{a}$ which yields a payoff lower than that of the SCF for every agent.

Several comments are worth pointing out about regime $\tilde{B}^*$. First, why do we employ a two-stage mechanism $\tilde{b}^*$? This is because we want to find an equilibrium in which a deviation from truth-telling can be identified and subsequently punished. This can be done only after utilities are learned, via the choice of “flag”.

Second, the agents report an integer also in the second stage of mechanism $\tilde{b}^*$. Note that either a positive integer or a flag leads to shutdown of strategic play in the regime.

This means that if we let the integer play to occur before the choice of flag an agent may deviate from truth-telling by reporting a false type and a positive integer, thereby avoiding subsequent detection and punishment altogether.

Third, note that the initial mechanism enforces an arbitrary outcome and only integers are reported. The integer play affects transitions such that the agents' continuation payoffs are bounded. We do not however allow for any strategic play towards first period implementation in order to avoid potential incentive or coordination problems; otherwise, one equilibrium may be that all players flag in period 1.

Next Lemma shows that Bayesian equilibria of regime $\tilde{B}^*$ exhibit the same payoff properties as those of regime B^* in Section 4.2 above and induce expected continuation payoff profile $v(f)$. This characterization result is obtained by applying similar arguments to those leading to Lemma 9 above. The additional complication in deriving the result for $\tilde{B}^*$ is that we also need to ensure that no agent flags on an equilibrium path. This is achieved inductively by showing that expected continuation payoffs are $v(f)$ and, hence, efficient, whereas flagging induces inefficiency because the continuation payoff for each i after flagging is $v_i(\tilde{a}) < v_i(f)$ (see the Appendix for the proof).

Lemma 10 *If f is efficient and satisfies conditions ω'' and v , every $\sigma \in Q^\delta(\tilde{B}^*)$ is such that $E_{\theta(t)}\pi_i^{t+1}(\sigma, B^*) = v_i(f)$ for any i , t and $\theta(t)$.*

We next establish that, with an identifiable SCF, $\tilde{B}^*$ admits a Bayesian Nash equilibrium that attains the desired outcome path with sufficiently patient agents.

Lemma 11 *Suppose that f satisfies identifiability and condition ω'' . If δ is sufficiently large, there exists $\sigma^* \in \overline{Q}^\delta(\tilde{B}^*)$ such that, for any $t > 1$, $\theta(t)$ and θ^t , (i) $g^{\theta(t)}(\sigma^*, \tilde{B}^*) = \tilde{b}^*$; and (ii) $a^{\theta(t), \theta^t}(\sigma^*, \tilde{B}^*) = f(\theta^t)$.*

Proof. By condition ω'' there exists some $\epsilon > 0$ such that, for each i , $v_i(\tilde{a}) < v_i(f) - \epsilon$. Define $\rho \equiv \max_{i, \theta, a, a'} [u_i(a, \theta) - u_i(a', \theta)]$ and $\bar{\delta} \equiv \frac{\rho}{\rho + \epsilon}$. Fix any $\delta \in (\bar{\delta}, 1)$. Consider the following symmetric strategy profile $\sigma^* \in \Sigma$: for each i , σ_i^* is such that:

- for any θ_i , $\sigma_i^*(\emptyset, Z, \theta_i) = 0$;
- for any $t > 1$ and corresponding history, if $\tilde{b}^*$ is played in the period,
 - in Stage 1, it always reports the true type;

- in Stage 2, it reports NF and zero integer if the agent has not detected a false report from another agent or has not made a false report himself in Stage 1; otherwise, report F .

From these strategies, each agent i obtains continuation payoff $v_i(f)$ at the beginning of each period $t > 1$. Let us now examine deviation by any agent i . First, consider $t = 1$. Given the definition of mechanism Z and transition rule 1(b), announcing a positive integer alters neither the current period's outcome/payoff nor the continuation payoff at the next period. Second, consider any $t > 1$ and any corresponding history on the equilibrium path. Deviation can take place in two stages:

(i) Stage 1 - Announce a false type. But then, due to identifiability, another agent will raise “flag” in Stage 2, thereby activating permanent implementation of $\tilde{a}$ as of the next period. The corresponding continuation payoff cannot exceed $(1 - \delta) \max_{a, \theta} u_i(a, \theta) + \delta(v_i(f) - \epsilon)$, while the equilibrium payoff is at least $(1 - \delta) \min_{a, \theta} u_i(a, \theta) + \delta v_i(f)$. Since $\delta > \bar{\delta}$, the latter exceeds the former, and the deviation is not profitable.

(ii) Stage 2 - Flag or announce a non-zero integer following a stage 1 at which a no agent provides a false report. But given transition rules 2(b) and 2(c), such deviations cannot make i better off than no “flag” and zero integer. ■

Notice from the proof that the above equilibrium strategy profile is also sequentially rational; in particular, in the Stage 2 continuation game following a false report (off-the-equilibrium), it is indeed a best response by either the liar or the detector to “flag” given that there will be another agent also doing the same. Furthermore, note that the mutual optimality of this behavior is supported by an indifference argument; outcome $\tilde{a}$ is permanently implemented regardless of one's response to another flag. This is in fact a simplification. Since $v_i(\tilde{a}) < v_i(f)$ for all i , for instance, the following modification to the regime makes it strictly optimal for an agent to flag given that there is another flag: if there is only one flag then the continuation regime simply implements $\tilde{a}$ forever, while if two or more agents flag the continuation regime alternates between enforcement of $\tilde{a}$ and dictatorships of those who flag such that those agents obtain payoffs greater than what they would obtain from $\tilde{a}$ forever (but less than the payoffs from f).

Lemmas 10 and 11 immediately enable us to state the following.

Theorem 5 *Consider the case of interdependent values. Suppose that f satisfies efficiency, identifiability, and conditions ω'' and v . Then, we have the following: there exist*

a regime R and $\bar{\delta} \in (0, 1)$ such that, for any $\delta \in (\bar{\delta}, 1)$ (i) the set $\bar{Q}^\delta(\tilde{B}^*)$ is non-empty and (ii) every $\sigma \in \bar{Q}^\delta(\tilde{B}^*)$ is such that $E\pi_i^t(\sigma, R) = v_i(f)$ for any i and for every $t \geq 2$.

The above result establishes payoff implementation from period 2 when $\delta \in (\bar{\delta}, 1)$. By the same reasoning as in the proof of Corollary 2, Theorem 5 can be extended to show outcome implementation if in addition f satisfies condition χ .

Private values In order to use the incentives of repeated play to overcome incentive compatibility, someone in the group must be able to detect a deviation and subsequently enforce punishment. With interdependent values, this is possible once utilities are learned; with private values, each agent's utility depends only on his own type and hence identifiability cannot hold.

One way to identify a deviation in the private value setup is to observe the distribution of an agent's type reports over a long horizon. By a law of large numbers, the distribution of the actual type realizations of an agent must approach the true prior distribution as the horizon grows. Thus, at any given history, if an agent has made type announcements that differ too much from the true distribution, it is highly likely that there have been lies, and it may be possible to build punishments accordingly such that the desired outcome path is supported in equilibrium. Similar methods, based on *review strategies*, have been proposed to derive a number of positive results in repeated games (see Chapter 11 of Mailath and Samuelson [16] and the references therein). Extending these techniques to our setup may lead to a fruitful outcome.

The budgeted mechanism of Jackson and Sonnenschein [12] in some sense adopts a similar approach. However, they use budgeting to derive a characterization of equilibrium properties that every equilibrium payoff profile must be arbitrarily close to the efficient target profile if the discount factor is sufficiently large and the horizon long enough.¹⁸ Their existence arguments on the other hand is not based on a budgeting (review strategy) type argument; in their setup the game played by the agents is finite and existence is obtained by appealing to standard existence results for a mixed equilibrium in a finite game. In our infinitely repeated setup we cannot appeal to such results because any game induced by any non-trivial regime will have infinite number of strategies.

¹⁸Recall that our approach delivers a sharper equilibrium characterization (i.e. precise, rather than approximate, matching between equilibrium and target payoffs obtained independently of δ) for a general incomplete information environment.

4.4 Ex post repeated implementation

The concept of Bayesian Nash equilibrium has drawn frequent criticism because it requires each agent to behave optimally against others for a given distribution of types and, hence, is sensitive to the precise details of the agents' beliefs about the others. We now show that our main results in this section can in fact be made sharper by addressing such a concern. To see how our Bayesian results above can be extended in this way, recall Lemmas 6 and 8 above and their proofs. These claims are established by deviation arguments that do not actually depend on the deviator's knowledge of the others' private information.

One way to formalize this point is to adopt the concept of *ex post equilibrium*, which requires the (private) strategy of each agent to be optimal against other agents' strategies for *every* possible realization of types. Though a weaker concept than dominant strategy implementation, ex post implementation in the one-shot setup is still an excessively demanding task (see Jehiel *at al* [13] and others). It turns out that, in our repeated setup, the same regimes used to obtain Bayesian implementation of efficient SCFs can deliver an even sharper set of positive results if the requirements of ex post equilibrium are introduced. The results below are obtained with the same set of conditions as in the complete information case, without the extra conditions v and χ assumed previously in this section for Bayesian repeated implementation.

In our repeated setup, informational asymmetry at each decision node concerns the agents' types in the past and present; the future is equally uncertain to all. Thus, for our repeated setup, it is natural to require the agents' strategies to be mutually optimal given any information that is available to all agents at every possible decision node of a regime (see Bergemann and Morris [4] for a definition of one-shot ex post implementation). To this end, denote, for any regime R and any strategy profile σ , each agent i 's expected continuation payoff of i at period $\tau \geq t$ conditional on knowing type realizations $(\theta(t), \theta^t)$ as follows: for $\tau = t$,

$$\pi_i^t(\sigma, R|\theta(t), \theta^t) = (1-\delta) \left[u_i(a^{\theta(t), \theta^t}, \theta^t) + \sum_{s \geq t+1} \sum_{\theta(s-t)} \sum_{\theta^s} \delta^{s-t} q(\theta(s-t), \theta^s) u_i(a^{\theta(t), \theta^t, \theta(s-t), \theta^s}, \theta^s) \right],$$

and, for any $\tau > t$,

$$\pi_i^\tau(\sigma, R|\theta(t), \theta^t) = (1-\delta) \sum_{s \geq \tau} \sum_{\theta(s-t)} \sum_{\theta^s} \delta^{s-\tau} q(\theta(s-t), \theta^s) u_i(a^{\theta(t), \theta^t, \theta(s-t), \theta^s}, \theta^s).$$

We shall then say that a strategy profile σ is an ex post equilibrium of R if, for any i , t , $\theta(t)$ and θ^t , $\pi_i^t(\sigma, R|\theta(t), \theta^t) \geq \pi_i^t(\sigma'_i, \sigma_{-i}, R|\theta(t), \theta^t)$ for all $\sigma'_i \in \Sigma_i$. Thus, at any history along the equilibrium path, the agents' strategies must be mutual best responses for every possible type realizations up to, and including, the period.¹⁹

Let $\overline{Q}^\delta(R)$ denote the set of ex post equilibria of regime R with discount factor δ . The definitions of repeated implementation with ex post equilibrium are analogous to those with Bayesian equilibrium as in Definition 4. We obtain the following result by considering regime B^* defined in Section 4.2. Here, the existence of an ex post equilibrium is ensured by *ex post incentive compatibility*, that is, $u_i(f(\theta), \theta) \geq u_i(f(\theta'_i, \theta_{-i}), \theta)$ for all i , $\theta = (\theta_i, \theta_{-i})$ and θ'_i .

Theorem 6 *Fix any $I \geq 2$. If f is efficient (in the range), ex post incentive compatible and satisfies condition ω (ω'), f is payoff-repeated-implementable in ex post equilibrium from period 2. If, in addition, f is strictly efficient (in the range), f is repeated-implementable in ex post equilibrium from period 2.*

The proof proceeds in a similar way to the arguments appearing in Section 3 for the complete information setup. It is therefore relegated to the Appendix. Note that Theorem 6 presents results that are, compared to the Bayesian results, sharper not only in terms of payoff implementation but also in terms of outcome implementation.

Our final result demonstrates how the same regime, $\tilde{B}^*$, that was used to drop incentive compatibility for Bayesian repeated implementation in the interdependent value case can also deliver analogous results for ex post repeated implementation. If the agents are sufficiently patient it is indeed possible to obtain the results without ex post incentive compatibility, which has been identified as incompatible with ex post implementation in the one-shot setup (Jehiel *et al* [13]).

¹⁹A stronger definition of ex post equilibrium would be to require that the strategies be mutually optimal at every possible *infinite* sequence of type profiles. Such a definition, however, would not be in the spirit of repeated settings in which decisions are made before future realizations of uncertainty. Also, such a definition would be excessively demanding. For example, in regime B^* ex post incentive compatibility no longer guarantees existence: if others always tell the truth and report zero integer, for some infinite sequence of states it may be better for a player to report a positive integer and become the odd-one-out.

Theorem 7 *Consider the case of interdependent values. Suppose that f satisfies efficiency, identifiability and condition ω'' . Then, we have the following: there exist a regime R and $\bar{\delta} \in (0, 1)$ such that, for any $\delta \in (\bar{\delta}, 1)$, (i) $\overline{Q}^\delta(R)$ is non-empty; and (ii) every $\sigma \in \overline{Q}^\delta(R)$ is such that $E\pi_i^t(\sigma, R) = v_i(f)$ for all i and $t \geq 2$.*

The proof appears in the Appendix. It adapts the arguments for Theorem 6 to those for Theorem 5, and notes that the strategy profile constructed to prove Lemma 11 also constitutes an ex post equilibrium of regime $\tilde{B}^*$.

5 Concluding discussion

5.1 Period 1

Our sufficiency Theorems 2-7 do not guarantee period 1 implementation of the SCF. We offer two comments. First, period 1 can be thought of as a pre-play round that takes place before the first state is realized. If such a round was available, one could ask the agents simply to announce a non-negative integer with the same transitions and continuation regimes such that each agent's equilibrium payoff at the beginning of the game corresponds exactly to the target level. The continuation regimes would then ensure that the continuation payoffs remained correct at every subsequent history.

Second, period 1 implementation can also be achieved if the players have a preference for simpler strategies at the margin, in a similar way that complexity-based equilibrium refinements have yielded sharper predictions in various dynamic game settings (see Chatterjee and Sabourian [8] for a recent survey). This literature offers a number of different definitions of complexity of a strategy and of equilibrium concepts in the presence of complexity.²⁰ To obtain our repeated implementation results from period 1, only a minimal addition of complexity aversion is needed. We need (i) *any* measure of complexity of a strategy under which a Markov strategy is simpler than a non-Markov strategy and (ii) a refinement of (Bayesian) Nash equilibrium in which each player cares for complexity of his strategy *lexicographically* after payoffs, that is, each player only chooses among the minimally complex best responses to the others' strategies. One can then show that every

²⁰A most well-known measure of strategic complexity is the size of minimal automaton (in terms of number of states) that implement the strategy (e.g. Abreu and Rubinstein [1]).

equilibrium in the canonical regimes above, with both complete and incomplete information, must be Markov and hence the main results extend to implementation from outset. We refer the reader to the Supplementary Material for formal statements and proofs.

5.2 Non-exclusive SCF

In our analysis thus far, repeated implementation of an efficient SCF has been obtained with an auxiliary condition ω (or its variation ω') which assumes that, for each agent, the (one-period) expected utility from implementation of the SCF must be bounded below by that of some constant SCF. The role of this condition is to construct, for each agent, a history-independent and non-strategic continuation regime in which the agent derives a payoff equal to the target level. We next define another condition that can fulfil the same role: an SCF is *non-exclusive* if, for each i , there exists some $j \neq i$ such that $v_i(f) \geq v_i^j$. The name of this property comes from the fact that, otherwise, there must exist some agent i such that $v_i(f) < v_i^j$ for all $j \neq i$; in other words, there exists an agent who strictly prefers a dictatorship by any other agent to the SCF itself.

Non-exclusion enables us to build for any i a history-independent regime that appropriately alternates dictatorial mechanisms $d(i)$ and $d(j)$ for some $j \neq i$ (instead of $d(i)$ and some trivial mechanism $\phi(a)$) such that agent i can guarantee payoff $v_i(f)$ exactly. We cannot say that either condition ω (or ω') or non-exclusion is a weaker requirement than the other.²¹

5.3 Off the equilibrium

In one-shot implementation, it has been shown that one can improve the range of achievable objectives by employing extensive form mechanisms together with refinements of Nash equilibrium as solution concept (Moore and Repullo [23] and Abreu and Sen [2] with complete information, and Bergin and Sen [6] with incomplete information, among others). Although this paper also considers a dynamic setup, the solution concept adopted is that of (Bayesian) Nash equilibrium and our characterization results do not rely on imposing a particular assumption off-the-equilibrium behavior or beliefs to “kill off” unwanted equilibria. At the same time, our existence results do not involve construction of Nash equilibria based on non-credible threats and/or unreasonable beliefs off-the-equilibrium.

²¹See the Supplementary Material for related examples.

Thus, we can replicate the same set of results with subgame perfect or sequential equilibrium as the solution concept.

A related issue is that of efficiency of off-the-equilibrium paths. In one-shot extensive form implementation, it is often the case that off-the-equilibrium inefficiency is imposed in order to sustain desired outcomes on-the-equilibrium. Several authors have, therefore, investigated to what extent the possibility of *renegotiation* affects the scope of implementability (for example, Maskin and Moore [19]). For many of our repeated implementation results, this needs not be a cause for concern since off-the-equilibrium outcomes in our regimes can actually be made efficient. Recall that the requirement of condition ω is that, for each agent i , there exists some outcome $\tilde{a}^i$ which gives the agent an expected utility less than or equal to that of the SCF. If the environment is rich enough, such an outcome can indeed be found on the efficient frontier itself. Moreover, if the SCF is non-exclusive, the regimes can be constructed so that off-the-equilibrium is entirely associated with dictatorial outcomes, which are efficient.

5.4 Other assumptions

In this paper, we have restricted our attention to implementation in pure strategies only. Mixed/behavior strategies can be incorporated into our setup. To see this, notice that the additional uncertainty arising from randomization does not alter the substance of deviation argument that ensures each agent's continuation payoffs in the canonical regimes to be bounded below exactly by the target payoff (see the proofs of Lemmas 2 and 6). Even if other agents randomize over an infinite support, an agent can guarantee that he will have the highest integer with probability arbitrarily close to 1.

Finally, this paper considers the case in which preferences follow an i.i.d. process. A potentially important extension is to generalize the process with which individual preferences evolve. However, since our definition of efficiency depends on the prior (ex ante) distribution, allowing for history-dependent distribution makes such efficiency a less natural social objective than in the present i.i.d. setup. Also, this extension will introduce the additional issue of learning. We shall leave these questions for future research.

6 Appendix

Proof of Theorem 3 Consider any SCF f that satisfies ω . Define mechanism $\hat{g} = (M, \psi)$ as follows: $M_i = \Theta \times \mathbb{Z}_+$ for all i and ψ is such that

1. if $m_1 = (\theta, \cdot)$ and $m_2 = (\theta, \cdot)$, then $\psi(m) = f(\theta)$;
2. if $m_1 = (\theta^1, \cdot)$ and $m_2 = (\theta^2, \cdot)$, and $\theta^1 \neq \theta^2$, then $\psi(m) \in L_1(\theta^2) \cap L_2(\theta^1)$ (by self-selection, this is well defined).

Regime $\hat{R}$ is defined to be such that $\hat{R}(\emptyset) = \hat{g}$ and, for any $h = ((g^1, m^1), \dots, (g^{t-1}, m^{t-1})) \in H^t$ such that $t > 1$ and $g^{t-1} = \hat{g}$:

1. if $m_1^{t-1} = (\cdot, 0)$ and $m_j^{t-1} = (\cdot, 0)$, then $\hat{R}(h) = \hat{g}$;
2. if $m_i^{t-1} = (\cdot, z^i)$, $m_j^{t-1} = (\cdot, 0)$ and $z^i \neq 0$, then $\hat{R}|h = S^i$, where $S^i \in \mathcal{S}(i, \tilde{a}^i)$ such that $v_i(f) \geq v_i(\tilde{a}^i)$ and $\pi_i(S^i) = v_i(f)$;
3. if m^{t-1} is of any other type and i is lowest-indexed agent among those who announce the highest integer, then $\hat{R}|h = D^i$.

We prove the claim via the following claims.

Claim 1: Fix any $\sigma \in \Omega^\delta(\hat{R})$. For any $t > 1$ and $\theta(t)$, if $g^{\theta(t)} = \hat{g}$, $\pi_i^{\theta(t)} \geq v_i(f)$.

This can be established by analogous reasoning to that behind Lemma 2.

Claim 2: Fix any $\sigma \in \Omega^\delta(\hat{R})$. Also, assume that, for each i , outcome $\tilde{a}^i \in A$ used in the construction of S^i above satisfies (4) in the main text. Then, for any t and $\theta(t)$, if $g^{\theta(t)} = \hat{g}$ then $m_i^{\theta(t), \theta^t} = (\cdot, 0)$ and $m_j^{\theta(t), \theta^t} = (\cdot, 0)$ for any θ^t .

Suppose not; then, for some t , $\theta(t)$ and θ^t , $g^{\theta(t)} = \hat{g}$ and the continuation regime next period at $\mathbf{h}(\theta(t), \theta^t)$ is either D^i or S^i for some i . By a similar reasoning as with the three-or-more-player case, it then follows that, for $j \neq i$,

$$\pi_j^{\theta(t), \theta^t} < v_j^j \quad (8)$$

Consider two possibilities. If the continuation regime is $S^i = \Phi(\tilde{a}^i)$ then $\pi_i^{\theta(t), \theta^t} = v_i(f) = v_i(\tilde{a}^i)$ and hence (8) follows from (4) in the main text. If the continuation regime

is D^i or $S^i \neq \Phi(\tilde{a}^i)$, $d(i)$ occurs in some period. But then (8) follows from $v_j(\tilde{a}^i) \leq v_j^j$ and $v_j^i < v_j^j$ (where the latter inequality follows from Assumption (A)).

Then, given (8), agent j can profitably deviate at $(\mathbf{h}(\theta(t)), \theta^t)$ by announcing the same state as σ_j and an integer higher than i 's integer choice at such a history. This is because the deviation does not alter the current outcome (given the definition of ψ of $\hat{g}$) but induces regime D^j in which, by (8), j obtains $v_j^j > \pi_j^{\theta(t), \theta^t}$. But this is a contradiction.

Claim 3: Assume that f is efficient in the range and, for each i , outcome $\tilde{a}^i \in A$ used in the construction of S^i above satisfies (4) in the main text. Then, for any $\sigma \in \Omega^\delta(\hat{R})$, $\pi_i^{\theta(t)} = v_i(f)$ for any i , $t > 1$ and $\theta(t)$.

Given Claims 1-2, and since f is efficient in the range, we can directly apply the proof of Lemma 4 in the main text.

Claim 4: $\Omega^\delta(\hat{R})$ is non-empty if self-selection holds.

Consider a symmetric Markov strategy profile in which, for any θ , each agent reports $(\theta, 0)$. Given ψ and self-selection, any unilateral deviation by i at any θ results either in no change in the current period outcome (if he does not change his announced state) or it results in current period outcome belonging to $L_i(\theta)$. Also, given the transition rules, a deviation does not improve continuation payoff at the next period either. Therefore, given self-selection, it does not pay i to deviate from his strategy.

Finally, given Claims 3-4, the proof of Theorem 3 follows by exactly the same reasoning as those of Theorem 2 and its Corollary.

Proof of Lemma 10 Fix any $\sigma \in Q^\delta(\tilde{B}^*)$. We proceed with the following claims.

Claim 1: Suppose that f is efficient and satisfies condition ω'' . Fix any $t > 1$. Assume that, for any $\theta(t)$, $g^{\theta(t)} = \tilde{b}^$ and also that, at any θ^t , every agent reports “no flag” in Stage 2. Then, for any i , $E_{\theta(t)}\pi_i^{t+1} = v_i(f)$ for all $\theta(t)$ and hence $E\pi_i^{t+1} = v_i(f)$.*

Since, in the above claim, it is assumed that every agent reports “no flag” in Stage 2 at any θ^t , the claim can be proved analogously to Lemmas 6 and 7 above.

Claim 2: Suppose that f is efficient and satisfies conditions ω'' and v . Fix any $t > 1$. Assume that, for any $\theta(t)$, $g^{\theta(t)} = \tilde{b}^$ and also that, at any θ^t , every agent reports “no flag” in Stage 2. Then, for any $\theta(t+1)$, the following two properties hold: (i) $g^{\theta(t+1)} = \tilde{b}^*$ and (ii) every agent will report “no flag” at period $t+1$ for any θ^{t+1} .*

Again, since in the above claim it is assumed that every agent reports “no flag” in Stage 2 at any θ^t , part (i) of the claim can be established analogously to Lemma 8 above.

To prove part (ii), suppose otherwise; then some agent reports “flag” in Stage 2 of period $t + 1$ at some sequence of states $\hat{\theta}^1, \dots, \hat{\theta}^{t+1}$.

For any s and $\theta(s) \in \Theta^{s-1}$ define SCF $f^{\theta(s)}$ by $f^{\theta(s)}(\theta) = a^{\theta(s), \theta}(\sigma, \tilde{B}^*)$ for all $\theta \in \Theta$. Then $E\pi_i^{t+1}$ can be written as

$$E\pi_i^{t+1} = (1 - \delta) \left[\sum_{\theta(t+1)} q(\theta(t+1)) v_i(f^{\theta(t+1)}) + \delta \sum_{s \geq t+2} \sum_{\theta(s)} \delta^{s-t-2} q(\theta(s)) v_i(f^{\theta(s)}) \right]. \quad (9)$$

Next, for any $s > t + 1$, let $\hat{\Theta}^{s-1} = \left\{ \theta(s) = (\theta^1, \dots, \theta^{s-1}) : \theta^\tau = \hat{\theta}^\tau \ \forall \tau \leq t + 1 \right\}$ be the set of $(s - 1)$ states that are consistent with $\hat{\theta}^1, \dots, \hat{\theta}^{t+1}$. Since there is flagging in period $t + 1$ after the sequence of states $\hat{\theta}^1, \dots, \hat{\theta}^{t+1}$ it follows that, for all i ,

$$(1 - \delta) \sum_{s \geq t+2} \sum_{\theta(s) \in \hat{\Theta}^{s-1}} \delta^{s-t-2} q(\theta(s)) v_i(f^{\theta(s)}) = \alpha v_i(\tilde{a}), \quad (10)$$

where $\alpha = (1 - \delta) \sum_{s \geq t+2} \sum_{\theta(s) \in \hat{\Theta}^{s-1}} \delta^{s-t-2} q(\theta(s))$. Also, by the previous claim, $E\pi_i^{t+1} = v_i(f)$ for all i . Therefore, it follows from (9), (10), and $v(\tilde{a}) \ll v(f)$ that, for all i ,

$$v_i(f) < \frac{(1 - \delta)}{(1 - \delta\alpha)} \left\{ \sum_{\theta(t+1) \in \Theta^t} q(\theta(t+1)) v_i(f^{\theta(t+1)}) + \delta \sum_{s \geq t+2} \sum_{\theta(s) \notin \hat{\Theta}^{s-1}} \delta^{s-t-2} q(\theta(s)) v_i(f^{\theta(s)}) \right\}. \quad (11)$$

Since $\sum_{\theta(t+1)} q(\theta(t+1)) = 1$ and $\alpha + (1 - \delta) \sum_{s \geq t+2} \sum_{\theta(s) \notin \hat{\Theta}^{s-1}} \delta^{s-t-2} q(\theta(s)) = 1$ by definition, it follows that

$$\sum_{\theta(t+1)} q(\theta(t+1)) + \delta \sum_{s \geq t+2} \sum_{\theta(s) \notin \hat{\Theta}^s} \delta^{s-t-2} q(\theta(s)) = \frac{(1 - \delta\alpha)}{(1 - \delta)}.$$

Therefore, by (11), f is not efficient; but this is a contradiction.

Claim 3: (i) $E\pi_i^2 = v_i(f)$ for all i ; (ii) $g^{\theta(2)} = \tilde{b}^*$ for any $\theta(2) \in \Theta$; and (iii) every agent will report “no flag” in Stage 2 of period 2 for any $\theta(2)$ and θ^2 .

Since $\tilde{B}^*(\emptyset) = Z$, we can establish (i) and (ii) by applying similar arguments as those in Lemmas 6-8 to period 1. Also, by similar reasoning for Claim 2 just above, it must be

that no player flags in period 2 at any $(\theta(2), \theta^2)$; otherwise the continuation payoff profile will be $v(\tilde{a})$ from period 3 and this is inconsistent with $E\pi_i^2 = v_i(f) > v_i(\tilde{a})$ for all i .

Finally, to complete the proof of Lemma 10, note that, by induction, it follows from Claims 1-3 that, for any t , $\theta(t)$ and θ^t , $g^{\theta(t)} = \tilde{b}^*$ and, moreover, every agent will report “no flag” in Stage 2 of period t . But then, Claims 1 and 3 imply that $E_{\theta(t)}\pi_i^{t+1} = v_i(f)$ for all i , t and $\theta(t)$.

Proof of Theorem 6 We shall first characterize the set of ex post equilibria of regime B^* constructed in Section 4.2. We proceed with the following claims. It is assumed throughout this proof that f satisfies condition ω and, for each i , outcome $\tilde{a}^i \in A$ used in the construction of S^i in regime B^* satisfies condition (4) in the main text.

Claim 1: Fix any $\sigma \in \overline{Q}^\delta(B^)$, t and $\theta(t)$, and suppose that $g^{\theta(t)}(\sigma, B^*) = b^*$. Then, $m_i^{\theta(t), \theta^t}(\sigma, B^*) = (\cdot, 0)$ for all i and θ^t .*

To show this, suppose otherwise; so, at some t and $\theta(t)$, $g^{\theta(t)} = b^*$ but $m_i^{\theta(t), \theta^t} = (\cdot, z^i)$ with $z^i > 0$ for some i and θ^t . Then there must exist some $j \neq i$ such that

$$\pi_j^{t+1}(\sigma, B^* | \theta(t), \theta^t) < v_j^j. \quad (12)$$

The reasoning for this is identical to that for Claim 2 in the proof of Lemma 3, except that $\pi_j^{\theta(t), \theta^t}$ there is replaced by $\pi_j^{t+1}(\sigma, B^* | \theta(t), \theta^t)$.

Next, consider j deviating to another strategy σ'_j which yields the same outcome path as the equilibrium strategy, σ_j , at every history, except at $(\mathbf{h}_j(\theta(t)), \theta_j^t)$, $\theta^t = (\theta_j^t, \theta_{-j}^t)$, where it announces the type announced by σ_j and an integer higher than any integer that can be reported by σ at this history.

Given ψ of mechanism b^* , such a deviation does not change the utility that j receives in period t under type profile θ^t ; furthermore, by the definition of B^* , the continuation regime is D^j and, hence, $\pi_j^{t+1}(\sigma'_i, \sigma_{-i}, B^* | \theta(t), \theta^t) = v_j^j$. But, given (12), this means that $\pi_j^t(\sigma'_i, \sigma_{-i}, B^* | \theta(t), \theta^t) > \pi_j^t(\sigma, B^* | \theta(t), \theta^t)$, which contradicts that σ is an ex post equilibrium of B^* .

Claim 2: Fix any $\sigma \in \overline{Q}^\delta(B^)$, t and $\theta(t)$. Then, (i) $g^{\theta(t)}(\sigma, B^*) = b^*$; and (ii) $a^{\theta(t), \theta^t}(\sigma, B^*) \in f(\Theta)$ for all θ^t .*

This follows immediately from Claim 1 above and the fact that $B^*(\emptyset) = b^*$.

Claim 3: Fix any $\sigma \in \overline{Q}^\delta(B^)$, t , $\theta(t)$ and θ^t . Then $\pi_i^{t+1}(\sigma, B^* \mid \theta(t), \theta^t) \geq v_i(f)$ for all i .*

Suppose not; so, for some $\sigma \in \overline{Q}^\delta(B^*)$, t , $\theta(t)$ and $\theta^t = (\theta_i^t, \theta_{-i}^t)$, $\pi_i^{t+1}(\sigma, B^* \mid \theta(t), \theta^t) < v_i(f)$ for some i . Then, consider i deviating to σ'_i identical to σ_i except at $(\mathbf{h}_i(\theta(t)), \theta_i^t)$ it announces the same type as σ_i but a positive integer. (By previous Claim, b^* is played and every agent announces zero at this history.)

By the definition of B^* and previous Claim, this deviation only changes the continuation regime to S^i from which on average $v_i(f)$ can be obtained. Thus, $\pi_i^t(\sigma'_i, \sigma_{-i}, B^* \mid \theta(t), \theta^t) > \pi_i^t(\sigma, B^* \mid \theta(t), \theta^t)$. This contradicts that σ is an ex post equilibrium of B^* .

Claim 4: Suppose that f is efficient in the range. Fix any $\sigma \in \overline{Q}^\delta(B^)$, t , $\theta(t)$ and θ^t . Then, for all i , $\pi_i^{t+1}(\sigma, B^* \mid \theta(t), \theta^t) = v_i(f)$.*

This follows from Claim 2-3 and efficiency in the range of f .

To complete the proof of the theorem, note first that the existence of an ex post equilibrium follows immediately from the equilibrium definition and ex post incentive compatibility of f . Next, note that payoff implementability then follows from the previous claim and noting that $E\pi_i^{t+1}(\sigma, B^*) = \sum_{\theta(t), \theta^t} q(\theta(t), \theta^t) \pi_i^{t+1}(\sigma, B^* \mid \theta(t), \theta^t)$ via analogous reasoning to that behind Theorem 2 and its Corollary. This reasoning also delivers the last part of the theorem on strict efficiency and outcome implementability.

Proof of Theorem 7 Consider regime $\tilde{B}^*$ in Section 4.3. It is straightforward to show that, given identifiability and condition ω'' , the same strategy profile appearing in the proof of Lemma 11 is also a ex post equilibrium of regime $\tilde{B}^*$ with sufficiently patient agents; this establishes part (i) of the theorem. To prove part (ii), it suffices to prove the following characterization of the ex post equilibria of regime $\tilde{B}^*$ (this result is the analogue to Lemma 10 that was used above to prove Bayesian implementation via $\tilde{B}^*$).

Claim: If f is efficient and satisfies condition ω'' , every $\sigma \in \overline{Q}^\delta(\tilde{B}^)$ is such that $\pi_i^{t+1}(\sigma, \tilde{B}^* \mid \theta(t), \theta^t) = v_i(f)$ for any i , t , $\theta(t)$ and θ^t .*

To prove this we proceed with the following steps that will lead to induction arguments.

Step 1: Fix any $\sigma \in \overline{Q}^\delta(\tilde{B}^*)$, t , $\theta(t)$ and θ^t , and suppose that $g^{\theta(t)}(\sigma, \tilde{B}^*) = \tilde{b}^*$. In addition, suppose that, in period t , every agent plays “no flag” in stage 2. Then, we have:

- (i) Every agent also reports zero integer and, hence, $g^{\theta(t), \theta^t} = \tilde{b}^*$.
- (ii) $\pi_i^{t+1}(\sigma, \tilde{B}^* \mid \theta(t), \theta^t) = v_i(f)$ for all i .

(iii) In period $t + 1$, every agent will play “no flag” at any θ^{t+1} .

Since condition ω'' subsumes both conditions ω and (4) in the main text, and since we assume that every agent plays “no flag” in period t , part (i) can be established analogously to Claim 1 in the proof of Theorem 6 above. Given efficiency of f , part (ii) can be proved using similar arguments to those behind Claims 3-4 in the same proof. (Here, note that we need efficiency over the entire set of SCFs, not just in the range. This is because we have had to make the additional assumption of “no flag”.)

To prove part (iii), suppose otherwise; so, at some θ^{t+1} , some agent plays “flag”. Given the transition rules, the continuation regime starting in $t + 2$ implements $\tilde{a}$ forever. Therefore, for all i , we have

$$\begin{aligned} \pi_i^{t+1}(\sigma, \tilde{B}^* | \theta(t), \theta^t) &= (1 - \delta) \sum_{\theta} p(\theta) u_i(a^{\theta(t), \theta^t, \theta}, \theta) \\ &\quad + \delta \sum_{\theta \in \Theta \setminus \theta^{t+1}} p(\theta) \pi_i^{t+2}(\sigma, \tilde{B}^* | \theta(t), \theta^t, \theta) + \delta p(\theta^{t+1}) v_i(\tilde{a}). \end{aligned} \quad (13)$$

But, for all i , $\pi_i^{t+1}(\sigma, \tilde{B}^* | \theta(t), \theta^t) = v_i(f)$ and $v_i(\tilde{a}) < v_i(f)$. Thus, by (13), we have

$$(1 - \delta) \sum_{\theta} p(\theta) u_i(a^{\theta(t), \theta^t, \theta}, \theta) + \delta \sum_{\theta \in \Theta \setminus \theta^{t+1}} p(\theta) \pi_i^{t+2}(\sigma, \tilde{B}^* | \theta(t), \theta^t, \theta) > \frac{v_i(f)}{1 - \delta p(\theta^{t+1})}, \quad (14)$$

for all i . Since $[\sum_{\theta} p(\theta) u_i(a^{\theta(t), \theta^t, \theta}, \theta)]_{i \in I} \in V$, $[\pi_i^{t+2}(\sigma, \tilde{B}^* | \theta(t), \theta^t, \theta)]_{i \in I} \in co(V)$ for all θ and $(1 - \delta) + \delta \sum_{\theta \in \Theta \setminus \theta^{t+1}} p(\theta) = 1 - p(\theta^{t+1})$, it follows from (14) that f is not efficient. But this is a contradiction.

Step 2: Fix any $\sigma \in \overline{Q}^{\delta}(\tilde{B}^*)$ and θ^1 . Then, we have: (i) $g^{\theta^1} = \tilde{b}^*$, (ii) $\pi_i^2(\sigma, \tilde{B}^* | \theta^1) = v_i(f)$ for all i , and (iii) in period 2, every agent plays “no flag” at any θ^2 .

We can apply the arguments for Step 1 above to period 1 and derive Step 2. The claim then follows from applying Steps 1-2 inductively.

References

- [1] Abreu, D. and A. Rubinstein, “The Structure of Nash Equilibria in Repeated Games with Finite Automata,” *Econometrica*, 56 (1988), 1259-1282.
- [2] Abreu, D. and A. Sen, “Subgame Perfect Implementaton: A Necessary and Almost Sufficient Condition,” *Journal of Economic Theory*, 50 (1990), 285-299.

- [3] Abreu, D. and A. Sen, "Virtual Implementation in Nash Equilibrium," *Econometrica*, 59 (1991), 997-1021.
- [4] Bergemann, D. and S. Morris, "Ex Post Implementation," *Games and Economic Behavior*, 63 (2008), 527-566.
- [5] Bergemann, D. and S. Morris, "Robust Implementation in Direct Mechanisms," *Review of Economic Studies*, 76 (2009), 1175-1204.
- [6] Bergin, J. and A. Sen, "Extensive Form Implementation in Incomplete Information Environments," *Journal of Economic Theory*, 80 (1998), 222-256.
- [7] Chambers, C. P., "Virtual Repeated Implementation," *Economics Letters*, 83 (2004), 263-268.
- [8] Chatterjee, K. and H. Sabourian, "Game Theory and Strategic Complexity," in *Encyclopedia of Complexity and System Science*, ed. by R. A. Meyers, Springer (2009).
- [9] Dasgupta, P., P. Hammond and E. Maskin, "Implementation of Social Choice Rules: Some General Results on Incentive Compatibility," *Review of Economic Studies*, 46 (1979), 195-216.
- [10] Dutta, B. and A. Sen, "A Necessary and Sufficient Condition for Two-Person Nash Implementation," *Review of Economic Studies*, 58 (1991), 121-128.
- [11] Jackson, M. O., "A Crash Course in Implementation Theory," *Social Choice and Welfare*, 18 (2001), 655-708.
- [12] Jackson, M. O. and H. F. Sonnenschein, "Overcoming Incentive Constraints by Linking Decisions," *Econometrica*, 75 (2007), 241-257.
- [13] Jehiel, P., M. Meyer-ter-Vehn, B. Moldovanu and W. R. Zame, "The Limits of Ex Post Implementation," *Econometrica*, 74 (2006), 585-610.
- [14] Jehiel, P. and B. Moldovanu, "Strictly Efficient Design with Interdependent Valuations," *Econometrica*, 69 (2001), 1237-1259.
- [15] Kalai, E. and J. O. Ledyard, "Repeated Implementation," *Journal of Economic Theory*, 83 (1998), 308-317.

- [16] Mailath, G. J. and L. Samuelson, *Repeated Games and Reputations: Long-run Relationships*, Oxford University Press (2006).
- [17] Maskin, E., "Auctions and Privatization," in *Privatization*, ed. by H. Siebert, Institut für Weltwirtschaft an der Universität Kiel (1992).
- [18] Maskin, E., "Nash Equilibrium and Welfare Optimality," *Review of Economic Studies*, 66 (1999), 23-38.
- [19] Maskin, E. and J. Moore, "Implementation and Renegotiation," *Review of Economic Studies*, 66 (1999), 39-56.
- [20] Maskin, E. and T. Sjöström, "Implementation Theory," in *Handbook of Social Choice and Welfare*, Vol. 1, ed. by K. Arrow *et al*, North-Holland (2002).
- [21] Matsushima, H., "A New Approach to the Implementation Problem," *Journal of Economic Theory*, 45 (1988), 128-144.
- [22] Mezzetti, C., "Mechanism Design with Interdependent Valuations: Efficiency," *Econometrica*, 72 (2004), 1617-1626.
- [23] Moore, J. and R. Repullo, "Subgame Perfect Implementation," *Econometrica*, 56 (1988), 1191-1220.
- [24] Moore, J. and R. Repullo, "Nash Implementation: A Full Characterization," *Econometrica*, 58 (1990), 1083-1099.
- [25] Mueller, E. and M. Satterthwaite, "The Equivalence of Strong Positive Association and Strategy-proofness," *Journal of Economic Theory*, 14 (1977), 412-418.
- [26] Saijo, T., "On Constant Maskin Monotonic Social Choice Functions," *Journal of Economic Theory*, 42 (1987), 382-386.
- [27] Sorin, S., "On Repeated Games with Complete Information," *Mathematics of Operations Research*, 11 (1986), 147-160.
- [28] Serrano, R., "The Theory of Implementation of Social Choice Rules," *SIAM Review*, 46 (2004), 377-414.