

Repeated Competing Mechanisms

Sambuddha Ghosh

Shanghai U. of Fin. & Economics

Seungjin Han

McMaster University

April 2016

Abstract

This paper studies the repeated game where *multiple* principals compete by offering short-term contracts to multiple agents, whose types are iid over periods. Agents send private messages to principals about their purported types as well as the contracts offered by other principals. The repeated game with patient players is shown to be tractable in the following two senses. First, all social choice functions (mappings from type profiles to action profiles) that can be supported with arbitrarily complex mechanisms can be supported with direct mechanisms on path, and all deviations can be punished by offering mechanisms whose message space asks for a type and an action report. Second, lower bounds on the payoffs with respect to arbitrarily complicated mechanisms may be calculated explicitly by solving simple programming problems. Neither result holds in the one-shot game.

Contents

1	Introduction	3
2	A general model of competing mechanisms	6
3	Complete information	8
3.1	One-shot game: role of mechanisms	9
3.2	Repeated game with complete information	13
4	Incomplete information	15
4.1	One-shot game with incomplete information	15
4.2	Repeated game with incomplete information	16
4.2.1	Incentive compatibility	16
4.2.2	Lower bound on a player's equilibrium payoff	18
4.2.3	A folk theorem for competing mechanisms	20
4.3	Comparison: one-shot vs. repeated games	22
5	Discussion	25
5.1	Markov types	25
5.2	Observability	25
5.3	Two agents	26
6	Conclusion	27
A	Appendix	29
A.1	Proof of Theorem 3	29
A.2	Proof of Theorem 4	31

1 Introduction

Because buyers are looking for better deals in the market, they are informed about contracts or terms of trade offered by sellers in the market. A seller, who might not be aware about what competing sellers are doing, thus has an incentive to come up with a sophisticated trading scheme that make buyers reveal their market information, on which the seller's terms of trade can then depend. For a simple example we can turn to a Bertrand duopoly, with a common constant marginal cost c . In a one-shot pricing game between duopolists, the monopoly price cannot be supported in equilibrium because firms have an incentive to cut prices. However as soon as sellers can offer price-matching contracts, they can use agents to implicitly collude even in the one-shot setting and charge a price higher than a competitive level. Unlike the single-principal setting, restricting attention to direct mechanisms and allowing terms of trade depend to depend only on buyer's payoff types, entails loss of generality as shown in McAfee (1993) and Peck (1997).

The above immediately leads to a conceptual question: What class of mechanisms can we restrict attention to in such settings? For the one-shot competing mechanism game, Epstein and Peters (1999) identifies a class of universal mechanisms. However there is no simple description of this class as the message spaces resemble the universal type-space for hierarchies of belief: the best response to any mechanism may be a more complicated mechanism; these are referred to as 'complex mechanisms' although they include all mechanisms, even very simple ones that take the same action after every history. Yamashita (2010) offers a simplification by showing that the subclass of *recommendation mechanisms* is also sufficient to support all equilibrium allocations when there are at least three agents: Each principal asks agents to suggest which direct mechanism he should offer, and *commits* to offer the direct mechanism recommended by a majority of agents. Recommendation mechanisms also encode punishments in the following sense—If any principal deviates from offering the recommendation mechanism, agents recommend the use of a mechanism that punishes the deviating principal in the same period.¹ However, each principal's worst equilibrium payoff is defined with reference to the set of mechanisms that are allowed in the competing mechanism game, which cannot be

¹The significance of having three or more agents is that one agent's deviation from the common recommendation has no effect.

parsimoniously described; thus the set of equilibrium allocations is not identified in terms of primitives of the model. Furthermore, no algorithm to construct them is available — if one knew the complex mechanisms that punish a certain principal, it would be simple to construct the corresponding recommendation mechanism; but we are not aware of a simple way of finding the direct mechanisms that principals must offer.

Following Yamashita (2010), we study a competing mechanism game with at least two principals and at least three sellers. We assume that a principal cannot make his action decision explicitly contingent on others' contracts but can commit to any mapping from agents' messages. We then study the repeated version of the above game where each agent's private payoff type is repeatedly drawn independently across periods. Specifically, our timeline is as follows. At the beginning of every period each principal offers a mechanism that is seen by the agents only. Each agent privately observes her payoff type and the mechanisms; they send private messages to the principals, who execute the mechanisms they offered at the beginning of the period. For simplicity, mechanisms and actions are assumed to be observable at the end of the period.²

We study perfect Bayesian equilibrium of the repeated game where all players are patient. Our repeated game combines short-term mechanisms with long-term interaction. Many real-life competing mechanism games are of this nature: for example, buyers repeatedly rent a car from one of two competing rental enterprises; various suppliers deliver raw materials to firms who manufacture similar products. Contracts govern the buyer-seller interaction, while a longer term interaction is ongoing, both among sellers and across the two sides of the market.

To describe our results we start with complete information, where agents' payoff types are fixed and observable. We find that the set of equilibrium allocations that can be supported in the repeated game are the same as those that can be supported in the one-shot game in the sense that the lower-bound of equilibrium payoffs of any principal is the minmax value over mixed actions in both games. Thus more complicated contracts do not expand the set of the equilibrium allocations.

However, as we now describe, results differ significantly when agents possess private information. A major contribution of our work is to show that, even in this

²Actually only agents need to observe the actions and mechanisms; see the discussion in the last section.

setting, the repeated game with patient players is tractable in two ways. First, ‘simple’ mechanisms very close to direct mechanisms (DMs) can mete out all punishments. Second, the lower bound on a principal’s equilibrium payoff admits a tractable characterisation in terms of primitives. This paper is thus very much about the complexity of mechanisms needed, rather than the equilibrium payoffs supported, as in a repeated game.

Let us see why this simplification obtains when players are patient. It is helpful to note the two points at which mechanisms may need to be complex. First, principals may be tempted to offer more complicated mechanisms on the equilibrium path so as to garner information about the mechanisms offered by the others and outwit them. We shall see that this one is easy to deter when the game is repeated because we do not need to deter deviations contemporaneously. A simple example might illustrate this: in the one-shot game duopolists might need to match the deviating seller’s price upon buyers’ reports on the deviator’s price in order to support high prices because they cannot do anything if they find out at the end of the game that their competitor undercut them. However, if the game is repeated and all players are patient they do not need to use this; a low price by a competitor can simply be punished by a price war later. However this observation may not be enough to simplify mechanisms because there is a second source of complexity off the equilibrium path—a principal who deviates from the path can offer complex mechanisms after he deviates; some of these might give him a higher utility, for example by rewarding agents for withholding information that would make it easier for the other principals to punish him. This problem doesn’t arise in the duopoly pricing game because one action, setting a zero price, punishes the deviating principal irrespective of what he does or what the types are, but such a punitive action doesn’t always exist. The key step in the paper is to show how this complexity can be tackled.

The difficulty is simplifying off-path punishments in such settings. Suppose that principal 1 must punish principal 2 for having deviated. First, we instruct the agents to induce a fixed action from the deviating principal that depends only on the mechanism that 2 offers (and nothing else). If agents do not do so, they can be identified and punished in the repeated game. This means that in the repeated game the deviating principal cannot do any better by offering an arbitrary mechanism off the path following his deviation than he could by offering a single, possibly mixed, action. Second, when the deviating principal cannot use very complicated best

responses, simple mechanisms suffice to mete out punishments.

In addition, we show that a weaker notion of incentive compatibility can be applied in the repeated game. In contrast to the one-shot game, these two equilibrium properties allow us to express the lower-bound on a principal's equilibrium payoff in the repeated game in terms of model primitives: It is equal to his maxmin value over his action space and the other principals' incentive compatible direct mechanisms conditional on the principal's action. These two equilibrium properties make the lower bound is lower than that in the one-shot setting. The weaker notion of incentive compatibility can also be applied to the equilibrium allocations. Combining it with the reduced lower-bound implies that the repeated game supports more allocations in equilibrium. When players are patient, principals can support an allocation that yields principals' payoffs above their lower bounds by offering only DMs on the equilibrium path but action-reporting DMs (ADMs) off the path following a principal's deviation: ADMs ask agents to report what action they are inducing from the deviating principal, along with their types.

2 A general model of competing mechanisms

We first describe the underlying one-shot game. The sets of principals and agents³ are, respectively, $\mathcal{J} := \{1, \dots, J\}$ and $\mathcal{I} := \{J + 1, \dots, J + I\}$ with $J \geq 2$ and $I \geq 3$. Each agent i has a type θ_i drawn from a finite set Θ_i according to the known distribution μ_i ; the joint distribution μ on $\Theta := \times_i \Theta_i$ is the product of marginals μ_i . Each principal j makes a decision a_j (henceforth referred to as an action) from a finite⁴ set A_j ; a random action of principal j is denoted by $\alpha_j \in \mathcal{A}_j := \Delta A_j$. A profile of pure actions is $a = (a_1, \dots, a_J) \in A := \times_{j \in \mathcal{J}} A_j$, while a profile of random actions is $\alpha \in \mathcal{A} := \times_{j \in \mathcal{J}} \mathcal{A}_j$, and $\mathcal{A}_{-j} := \times_{k \neq j} \mathcal{A}_k$. The vN-M (von Neumann-Morgenstern) expected utility function for player ℓ (principal or agent) is $u_\ell: \mathcal{A} \times \Theta \rightarrow \mathbb{R}$; payoffs are uniformly bounded by $\bar{u} < \infty$, i.e. $|u_\ell(\alpha, \theta)| < \bar{u}$ for all $\alpha \in \mathcal{A}$, all $\ell \in \mathcal{I} \cup \mathcal{J}$,

³To avoid confusion, we use feminine pronouns for agents and masculine pronouns for principals.

⁴Finiteness of the type and action spaces is not critical for our results, but are usually made in the literature. With a modicum of technicalities we can deal with a compact set of actions and a countable type-space.

and all $\theta \in \Theta$. All this information is encapsulated in the *underlying game*:

$$G := (\mathcal{J}; \mathcal{I}; (\mathcal{A}_j)_{j \in \mathcal{J}}; (\Theta_i)_{i \in \mathcal{I}}; (\mu_i)_{i \in \mathcal{I}}; (u_\ell)_{\ell \in \mathcal{J} \cup \mathcal{I}}). \quad (1)$$

The underlying game can be thought of as the simplest game where each principal's strategy space is the set of random (mixed) actions.

Fix an underlying game G as in (1), a collection of compact sets $\{M_j = \times_{i \in \mathcal{I}} M_{ij} \mid j \in \mathcal{J}\}$, and a collection $\{\Gamma_j \mid j \in \mathcal{J}\}$ where Γ_j is the set of all continuous mappings from the domain M_j to $\mathcal{A}_j$; let $\Gamma := \times_j \Gamma_j$. A mechanism is sufficiently general in terms of the degree of communication so that the class of mechanisms contains all best-responses to any profile of mechanisms the others may choose. The one-shot competing mechanism game (G, Γ) is the game with the following timing of moves:

1. Each principal j simultaneously offers a mechanism γ_j from Γ_j .
2. After observing mechanisms, all agents simultaneously send private messages, one to each principal, without observing others' messages; agent i 's message to j is $m_{ij} \in M_{ij}$.
3. A principal's action is determined by his mechanism,⁵ given the messages he receives, so that principal j takes action $\gamma_j(m_j) \in \mathcal{A}_j$ when he receives the profile of messages $m_j := (m_{ij})_{i \in \mathcal{I}} \in M_j$.
4. Finally, each player $\ell \in \mathcal{I} \cup \mathcal{J}$ earns the payoff $u_\ell(\gamma_1(m_1), \dots, \gamma_J(m_J), \theta)$.

The following assumptions are maintained throughout: first, messages from agent i to principal j private, i.e. it is not observable by other players, principals or agents. However the mechanisms offered and the actions chosen are assumed to be revealed at the end of the period to principals and agents; while eventual observation of mechanisms by the other principals may be regarded as too strong, our results would go through substantively unchanged even if principals never observed the others' mechanisms (see Section 5.2).

We also allow all players to observe the draw of a public correlation device at the beginning of the period, before principals offer their mechanisms. The PCD

⁵The class mechanisms is general enough to include everyday mechanisms and recommendation mechanisms in Yamashita (2010).

makes correlated actions feasible. An allocation or social choice function is a mapping $f : \Theta \rightarrow \Delta A$ from type profiles to probability distributions over actions. A direct mechanism with outcome f under truthtelling can be formulated in the following ways. First, all players observe the draw ω of a random variable distributed uniformly on the interval $[0, 1)$ and principal j takes an action $\pi_j(\theta, \omega)$. In this formulation principals use a mechanism $\pi_j : \Theta \times [0, 1) \rightarrow A_j$, which is essentially a pure mechanism conditional on any draw. However when punishing a principal for a deviation we might want the other principals to play a truly (ie not a predictable) mixed action. This is why it is useful to think of a direct mechanism as a mapping $\pi_j : \Theta \times [0, 1) \times [0, 1) \rightarrow A_j$ where the second variable is privately observed. For notational convenience we assume that the first randomness (correlation) has already been resolved and write the second one (independent mixing) is yet to be resolved; accordingly our mechanism maps into the space of random actions; henceforth a direct mechanism is $\pi_j : \Theta \rightarrow \Delta A_j$; this notation, maintaining consistency with a large part of the literature, means that the principal's mixing probabilities are observed. Alternatively we could have the principals announcing the mixing probabilities in each period and then we could employ a statistical test to check if the actual distributions are close to the distributions announced; since the game is repeated and players are patient, standard laws of large numbers apply. However this issue seems orthogonal to the main concern of this work and we circumvent it by writing the mechanism as $\pi_j : \Theta \rightarrow \Delta A_j$ and assuming that mixed actions are observed. We employ PBE (perfect Bayesian equilibrium) as the equilibrium solution concept.

3 Complete information

In this section, we focus on complete information, i.e. agents do not have private information about types; we model this by letting μ_i be *degenerate* at the type profile $\theta = (\theta_{J+1}, \dots, \theta_{J+I}) \in \Theta = \times_{i \in \mathcal{I}} \Theta_i$, where θ_i is the fixed type of agent i . For economy of notation, we consider two principals ($J = 2$) in this section, but this is without loss of generality. The point in studying mechanism design without private information is to highlight the role of mechanisms in a world with multiple principals because each principal wants to extract information from agents about

the mechanisms offered by all other principals. In addition to types that represent exogenous private information, competition among principals generates *endogenous private information*: agents have better market information than principals. The complete-information model shows how competing mechanisms can be used to sustain certain choices (of type-contingent action) using messages from the agents.

3.1 One-shot game: role of mechanisms

As a baseline model, consider the one-shot game with complete information but no mechanisms, i.e. the underlying game G . Agents play no role; this is the standard one-shot game where each principal j independently takes an action $\alpha_j \in \mathcal{A}_j$. The set of PBE payoffs coincides with Nash equilibrium payoffs, say $N(G)$. With a public correlation device, we can get payoffs in the convex hull $co(N(G))$.

We shall show that even with complete information and without repetition, agents play a vital role as soon as we allow the principals to use mechanisms rather than confining them to simple actions; this creates additional equilibrium payoffs in the one-shot game that are not Nash payoffs of the underlying game. To characterize the set of equilibrium allocations and payoffs, we first derive the lower bound on each principal j 's equilibrium payoff using recommendation mechanisms, where each principal asks agents to recommend actions and commits to playing the action that is recommended by a majority of agents; if one principal unilaterally deviates from offering such a mechanism, all agents recommend that the other principals choose action profiles that punish the deviating principal. Thus these recommendation mechanisms, proposed in Yamashita (2010), not only implement the equilibrium actions but also encode punishments to punish deviating principals.

Under complete information, agents only need to recommend an action to each principal. Each principal ℓ offers a recommendation mechanism

$$r_\ell : \mathcal{A}_\ell^I \rightarrow \mathcal{A}_\ell,$$

which leads principal ℓ to pick α'_ℓ if the majority of messages is α'_ℓ .⁶

Let $\underline{u}_j(\gamma_\ell, \gamma_j, \theta)$ denote principal j 's payoff in the worst continuation equilibrium of the one-shot messaging game (among agents) after the vector of the mechanisms

⁶Recommendation mechanisms in Yamashita (2010) reduce to this under complete information.

(γ_ℓ, γ_j) is offered. The minmax value (or payoff) of principal j in the one-shot game (G, Γ) is

$$w_j^1 := \min_{\gamma_\ell \in \Gamma_\ell} \max_{\gamma_j \in \Gamma_j} u_j(\gamma_\ell, \gamma_j, \theta).$$

The next lemma shows this general minmax value w_j^1 equals the simple minmax value, where all principals are allowed to use only actions (or, equivalently, constant mechanisms). It also shows how to ‘threaten’ to punish principal j with the payoff w_j^1 using recommendation mechanisms.

Lemma 1 *The minmax value of any principal j in the complete-information competing mechanism game (G, Γ) equals the minmax of the underlying game G :*

$$w_j^1 = \min_{\alpha_\ell \in \mathcal{A}_\ell} \max_{\alpha_j \in \mathcal{A}_j} u_j(\alpha_\ell, \alpha_j, \theta). \quad (2)$$

Proof. First of all, note that a simple action α_j is a constant mechanism that always assigns α_j regardless of agents’ messages. If principal j deviates to a simple action α_j , agents can recommend to principal ℓ an action that minimizes principal j ’s payoff conditional on α_j , which is

$$\bar{\varphi}_\ell^j(\alpha_j) := \arg \min_{\alpha_\ell \in \mathcal{A}_\ell} u_j(\alpha_\ell, \alpha_j, \theta). \quad (3)$$

This is the worst punishment that principal ℓ can induce in a continuation equilibrium upon principal j playing α_j .⁷ The maximum payoff that principal j can receive by playing a simple action is then

$$\max_{\alpha_j \in \mathcal{A}_j} u_j(\bar{\varphi}_\ell^j(\alpha_j), \alpha_j, \theta) = \max_{\alpha_j \in \mathcal{A}_j} \min_{\alpha_\ell \in \mathcal{A}_\ell} u_j(\alpha_\ell, \alpha_j, \theta).$$

Since the class of mechanisms Γ_j includes the set of actions $\mathcal{A}_j$, we have

$$w_j^1 \geq \max_{\alpha_j \in \mathcal{A}_j} \min_{\alpha_\ell \in \mathcal{A}_\ell} u_j(\alpha_\ell, \alpha_j, \theta). \quad (4)$$

Suppose that principal l commits to taking the action that is recommended by a majority of agents; let agents recommend the action $\underline{\alpha}_\ell^j$ to principal ℓ , regardless

⁷No agent can change this by a unilateral deviation.

of principal j 's mechanism, where

$$\underline{\alpha}_\ell^j \in \arg \min_{\alpha_\ell \in \mathcal{A}_\ell} \left\{ \max_{\alpha_j \in \mathcal{A}_j} \underline{u}_j(\alpha_\ell, \alpha_j, \theta) \right\}; \quad (5)$$

the maximum payoff that principal j can receive is then

$$\min_{\alpha_\ell \in \mathcal{A}_\ell} \max_{\alpha_j \in \mathcal{A}_j} u_j(\alpha_\ell, \alpha_j, \theta).$$

Since this is one way of punishing principal j , we have

$$w_j^1 \leq \min_{\alpha_\ell \in \mathcal{A}_\ell} \max_{\alpha_j \in \mathcal{A}_j} u_j(\alpha_\ell, \alpha_j, \theta) \quad \forall j \in \mathcal{J}. \quad (6)$$

Since minmax and maxmin values are equal with random actions (von Neumann's Minmax Theorem), (4) and (6) immediately imply (2). ■

Lemma 1 expresses each principal's minmax payoff in terms of primitives.

A (possibly correlated) action profile $\alpha \in \Delta A$, which chooses the pure action profile $a \in A$ with probability $\alpha[a]$, is strictly individually rational (SIR) for principals if

$$u_j(\alpha, \theta) = \sum_{a \in A} u_j(a, \theta) \alpha[a] > w_j^1 \quad \forall j \in \mathcal{J},$$

where w_j^1 is the principal's one-shot minmax value given by equation (2). The set of (correlated) action profiles that induce SIR payoffs is

$$\mathcal{F}^* := \{\alpha \in \Delta A \mid u_j(\alpha, \theta) > w_j^1 \quad \forall j \in \mathcal{J}\}. \quad (7)$$

Let $\bar{\mathcal{F}}^*$ be the closure of $\mathcal{F}^*$, which includes vectors of payoffs in which some principal's payoff may be equal to his minmax value. The lower bounds identified in Lemma 1 characterize the set of equilibrium allocations as follows. The proof in Yamashita (2010) simplifies to this under complete information.

Theorem 1 *In the one-shot competing mechanism game (G, Γ) ,*

1. *(correlated) action profiles in $\bar{\mathcal{F}}^*$ are supportable in a PBE;*
2. *such a PBE can be constructed using only recommendation mechanisms (r_1, r_2) ,*

under which each principal j asks every agent to send a message in his space of actions $\mathcal{A}_j$ and commits to take the action recommended by a majority.

Proof. Any action profile that gives any principal j a payoff strictly below his own minmax value cannot be sustained in equilibrium so that any action profile is supportable in a PBE must be SIR and therefore cannot be outside $\bar{\mathcal{F}}^*$. All payoffs within this set are feasible thanks to the public correlation device. In equilibrium, each principal ℓ offers a recommendation mechanism $r_\ell : \mathcal{A}_j^I \rightarrow \mathcal{A}_j$ regardless of the realization of the public correlation device. Fix $\alpha^* \in \bar{\mathcal{F}}^*$. Given the realization of the public correlation device, suppose that $\alpha = (\alpha_1, \alpha_2) \in \mathcal{A}$ in the support of α^* is the action profile that needs to be supported in equilibrium. Agents follow three rules:

1. if principals offer $r = (r_1, r_2)$, all agents send the message α_j to principal j ;
2. if principal j unilaterally deviates and offer anything other than r_j , agents recommend the action profile $\underline{\alpha}_\ell^j$, defined in (5), to principal $\ell \neq j$, and send messages to principal j that form an equilibrium of the induced one-shot messaging game among the agents (naturally what messages are allowed depends on the mechanism j deviated to);
3. if anything else happens agents can send to non-deviating principals any messages that represent a continuation equilibrium of the induced one-shot messaging game among the agents.

Obviously, agents cannot deviate unilaterally and impact the equilibrium because $I \geq 3$. If principal j deviates, agents recommend that the other principal ℓ takes the action $\underline{\alpha}_\ell^j$, which by Lemma 1 gives j a payoff w_j^1 , which is no better than the equilibrium payoff $u_j(\alpha, \theta)$ by hypothesis. So it is a best response for principal j to offer r_j . ■

If random actions are allowed so that the maxmin and minmax are the same, then it significantly reduces the amount of the information that a non-deviator needs to know in order to punish a deviator. Because the non-deviating principal ℓ can simply take the action $\underline{\alpha}_\ell^j$ that minmaxes deviating principal j , he only needs to know whether or not the other principal has deviated (more generally, with more

than two principals he needs to know the identity of the deviator, if any). Therefore, each principal offers a recommendation mechanism in equilibrium, and leaves agents to recommend which action (equilibrium or punitive) he needs to take.

3.2 Repeated game with complete information

The infinitely repeated game $(G, \Gamma)^\infty(\delta)$ involves playing the competing mechanism stage-game G at each time $t \geq 1$, with a common discount factor $\delta \in (0, 1)$ across periods, where principals are allowed to use mechanism profiles in Γ . Let us fix the type profile at $\theta = (\theta_{J+1}, \dots, \theta_{J+I}) \in \Theta = \times_{i \in \mathcal{I}} \Theta_i$ for all periods. Let $\gamma^t \in \Gamma$ and $\alpha^t \in \mathcal{A}$ be the mechanisms offered and the actions chosen at time t by the principal. Starting with the null history h^{-1} , a t -period history h^t is constructed from the $(t-1)$ -period history according to the formula $h^t = h^{t-1} \circ (\gamma^t, \alpha^t)$, where $\circ$ denotes concatenation. At the end of period t , both agents and principals observe the history h^t .⁸ Enforcement is harder if *principals cannot observe* mechanisms offered or actions taken by the other principals; this is discussed later. The (average) discounted payoff of player $\ell \in \mathcal{J} \cup \mathcal{I}$ from period τ onwards is $(1 - \delta) \sum_{t \geq \tau} \delta^{t-\tau} u_\ell(\alpha^t, \theta)$, where α^t is the action profile taken at time t .

Recall that the set of correlated action profiles that induce feasible and strictly individually rational payoffs in the one-shot game (G, Γ) is $\mathcal{F}^* = \{\alpha^* \in \Delta A \mid u_j(\alpha) > 0 \in \text{supp}(\alpha^*) \text{ is SIR for principals}\}$. We showed that w_j^1 in the one-shot game (G, Γ) is the same as j 's minmax value over complex mechanisms in Γ . This argument is valid because even if a non-deviating principal cannot change his mechanism upon a competing principal's deviation, the recommendation mechanism can induce the same effect. In the repeated game, a non-deviating principal can actually change his mechanism after a competing principal's mechanism and hence the same w_j^1 is of course j 's minmax value, which shown to be equal to the minmax value over actions in Lemma 1.

Theorem 2 (Complete-information folk theorem) *Let (G, Γ) be any one-shot competing mechanism game with the complete information. Then there exists $\underline{\delta} < 1$ such that for any $\delta \geq \underline{\delta}$*

⁸Our notational convention for time-dependent variables is to use $\varkappa^t$ to denote the value at period t of the variable $\varkappa$ of the stage-game, with the understanding that t is a superscript and not an exponent.

1. *any (correlated) action profile in $\mathcal{F}^*$ is the outcome of a PBE of the infinitely repeated game $(G, \Gamma)^\infty(\delta)$;*
2. *such a PBE can be supported using simple actions, i.e., constant mechanisms on and off the path.*

This is nothing but the standard folk theorem for infinitely repeated games, with the minmax value given by Lemma 1. The set of equilibrium allocations supported in the one-shot game is $\bar{\mathcal{F}}^*$, the closure of $\mathcal{F}^*$, whereas the set of equilibrium allocations supported in Theorem 2 in the repeated game based on the folk theorem argument is $\mathcal{F}^*$. Because any equilibrium allocation in the one-shot game is also an equilibrium allocation in the repeated game, the set of equilibrium allocations that can be supported in the repeated game is also $\bar{\mathcal{F}}^*$.

The key message is that mechanisms that principals need to use on and off the path to support equilibrium allocations in $\mathcal{F}^*$ are simpler in the repeated game than in the one-shot game. In the latter it is essential for principals to commit themselves to take a certain action according to the rule described in the recommendation mechanism. However, in the repeated game, they do not need commitment power because principals only need to take single actions on and off the path—each principal ℓ only needs to take his equilibrium action on the path, and $\underline{\alpha}_\ell^j$ off the path following principal j 's deviation.⁹

Notably, in the repeated game, principals who make equilibrium actions correlated by choosing their actions contingent on the realization of the PCD. However, in the one-shot game, agents make equilibrium actions correlated by choosing the recommendation of an action to each principal contingent on the realization of the PCD given each principal's recommendation mechanism; each principal offers his recommendation mechanism independent of the realization of the PCD in the one-shot game.

⁹The folk theorem does not cover those allocations in $\bar{\mathcal{F}}^* \setminus \mathcal{F}^*$ that induce payoffs where some principals get exactly equal their minmax values. These allocations may be supported in equilibrium of the repeated game if principals offer recommendation mechanisms rather than just direct mechanisms.

4 Incomplete information

Let $\mathcal{F}$ be the set of all possible correlated stage-SCFs $f: \Theta \rightarrow \Delta A$. A stage-SCF can be effected under truthtelling if principals offer a probability distribution over direct mechanisms; a profile of direct mechanisms (DMs) $\pi = (\pi_1, \dots, \pi_J)$, one for each principal, where direct mechanism (DM) offered by principal j is denoted by a mapping $\pi_j: \Theta \rightarrow \mathcal{A}_j$. Let Π_j be the set of all DMs available to principal j , the nature of which depends on the application under consideration; let $\Pi := \times_{j \in \mathcal{J}} \Pi_j$. Recall that $J \geq 2$ and $I \geq 3$.

4.1 One-shot game with incomplete information

Yamashita (2010) studies equilibrium allocations supportable in the incomplete-information game of competing mechanisms. Let γ be a profile of mechanisms offered, and let $u_j^1(\gamma)$ denote the worst payoff of principal j over all equilibria of the messaging game among agents that is induced by γ . The minmax value of principal j is then defined as

$$w_j^1 := \min_{\gamma_j \in \Gamma_{-j}} \max_{\gamma_j \in \Gamma_j} u_j^1(\gamma_{-j}, \gamma_j). \quad (8)$$

Suppose that we wish to construct an equilibrium where the principals offers the profile of DMs π^* . This can be done if all principals offer recommendation mechanisms; the only difference with the complete information case is that agents now have to recommend direct mechanisms, which of course have to satisfy UIC. On the equilibrium path all agents recommend π_j^* to each principal j ; if any principal j unilaterally deviates all agents recommend the direct mechanisms to the others that give the lowest possible payoff to principal j in some (UIC) continuation equilibrium. While anything that can be supported with arbitrarily complex mechanisms can also be supported with recommendation mechanisms, the minmax value cannot be found in terms of model primitives; therefore it is not easy to determine which payoffs are indeed equilibrium payoffs. Although recommendation mechanisms feel more familiar because they involve recommending the ubiquitous direct mechanisms, this model is still hard to apply because we do not know how to construct the actual direct mechanisms that agents must recommend.

4.2 Repeated game with incomplete information

We now turn to the repeated version of the game of incomplete information, which is our main object of study. As before, the repeated game allows us to relax incentive compatibility because deviation can be deterred if they can be detected at the end of the period. We first explain the right notion of IC in this dynamic setting.

4.2.1 Incentive compatibility

The key reason why the repeated game is simpler is that we can relax incentive compatibility; we start by explaining this intuition, which comes through even in the case with one agent with two possible types, θ and θ' . Suppose that each principal j offers a DM π_j such that $\pi_j(\theta) = \alpha_j$ and $\pi_j(\theta') = \alpha'_j$. The agent sends a pair of messages, the first to principal 1 and the second to principal 2, and induces the following action profiles:

(θ, θ)	(θ, θ')	(θ', θ)	(θ', θ')
(α_1, α_2)	(α_1, α'_2)	(α'_1, α_2)	(α'_1, α'_2)

Pairs of type messages such as (θ, θ') and (θ', θ) are inconsistent, i.e. the agent's type message to one principal is different from her type message to the other. In the one-shot game, incentive compatibility should be defined over all four profiles of type messages because principals cannot punish the agent even if they can detect her lies at the end of the period; such a profile of mechanisms is called UIC (unrestrictedly incentive compatible). In the repeated game, if (α_1, α'_2) or (α'_1, α_2) is taken by principals at the end of the period, it shows that the agent lied to at least one principal; principals can punish the agent subsequently. But if (α_1, α_2) or (α'_1, α'_2) occurs, principals do not know whether the agent lied or not because her type messages are consistent (i.e., messages to both principals are the same); so we still need incentive compatibility over $\{(\theta, \theta), (\theta', \theta')\}$. Therefore, in the repeated game profiles of mechanisms do not need to impose incentive compatibility over message profiles that induce actions that could not arise if the agent reported truthfully. We say that such a profile of mechanisms is CIC (constrained incentive compatible).

Since our model has multiple agents, CIC has to be defined with more care because it is possible that some inconsistent message profile of agent i leads to to off-path actions only for certain consistent type messages by agents other than

i ; however detection with positive probability rather than certainty is enough to punish patient agents in future. Therefore we impose incentive compatibility over only profiles of messages, consistent or not, that necessarily induces an action profile that is on path, but not on a profile of inconsistent type messages that induces with positive probability an action profile that cannot happen under truthful reporting. Now let us formalize the notion of IC with multiple agents. Given a profile of DMs $\pi := (\pi_1, \dots, \pi_J)$ in a period, the expected payoff of agent i of type θ_i , when the other agents truthfully report their types, is

$$\mathbb{E}_{\mu_{-i}} [u_i(\pi_1(\theta_{i1}, \theta_{-i}), \dots, \pi_J(\theta_{iJ}, \theta_{-i}), \theta_i, \theta_{-i})],$$

where $\mathbb{E}_{\mu_{-i}} [\cdot]$ is the expectation operator with respect to the probability distribution μ_{-i} over Θ_{-i} .

Given any profile of DMs π , the set of all actions profiles induced by truthful type reports is given by

$$\hat{\mathcal{A}}(\pi) := \{\pi(\theta) \mid \theta \in \Theta\}, \text{ where } \pi(\theta) := (\pi_1(\theta), \dots, \pi_J(\theta)) \forall \theta \in \Theta.$$

For any given π , we define $B_i(\pi) \subset (\Theta_i)^J$ for each i as the set of all profiles of type reports of agent i , one report to each principal, that lead to an action profile in $\hat{\mathcal{A}}(\pi)$ irrespective of the types of the other agents as long as they report truthfully:

$$B_i(\pi) := \left\{ (\theta_{i1}, \dots, \theta_{iJ}) \in (\Theta_i)^J \mid [\pi_1(\theta_{i1}, \theta_{-i}), \dots, \pi_J(\theta_{iJ}, \theta_{-i})] \in \hat{\mathcal{A}}(\pi) \forall \theta_{-i} \in \Theta_{-i} \right\}.$$

We propose below two different definitions of incentive compatibility (IC), both defined on the agent's stage-game payoff, depending on what kind of misreports are being deterred.

Definition 1 *A profile of DMs $\pi = (\pi_1, \dots, \pi_J)$ satisfies IC over $D_i \subset (\Theta_i)^J$ with respect to (w.r.t.) μ if for all $i \in \mathcal{I}$ and all $\theta = (\theta_i, \theta_{-i}) \in \Theta$ we have*

$$\begin{aligned} \mathbb{E}_{\mu_{-i}} [u_i(\pi(\theta), \theta)] &\geq \\ \mathbb{E}_{\mu_{-i}} [u_i(\pi_1(\theta_{i1}, \theta_{-i}), \dots, \pi_J(\theta_{iJ}, \theta_{-i}), \theta)] &\quad \forall (\theta_{i1}, \dots, \theta_{iJ}) \in D_i, \end{aligned} \quad (9)$$

where $\pi(\theta) = [\pi_1(\theta_i, \theta_{-i}), \dots, \pi_J(\theta_i, \theta_{-i})]$.

1. If π satisfies (9) over $D_i = (\Theta_i)^J$ we say that π is unrestrictedly incentive compatible (UIC).
2. If π satisfies (9) over $D_i = B_i(\pi)$, we say that π is constrained incentive compatible (CIC). Clearly $(\theta_i, \dots, \theta_i) \in B_i(\pi)$ for any $\theta_i \in \Theta_i$.

Letting Π^U (respectively Π^C) denote the set of all possible profiles of UIC (CIC) DMs, $(\Theta_i)^J \supset B_i(\pi)$ clearly implies $\Pi^U \subset \Pi^C$: a weaker notion of IC needs to deter fewer deviations.

Let us clarify the notions of IC in terms of the DMs that belong to each class, and the misreports that each notion deters using *contemporaneous* incentives. Consider agent i contemplating a unilateral deviation from truth telling. We say that agent i reports *consistently* if she reports the same type to all principals (i.e. $\tilde{\theta}_{i1} = \dots = \tilde{\theta}_{iJ}$), and *inconsistently* if $\{\tilde{\theta}_{i1}, \dots, \tilde{\theta}_{iJ}\}$ contains at least two distinct elements, i.e. she sends different messages to at least two principals. Consistent reports can be false: agent i sends $\theta'_i = \tilde{\theta}_{i1} = \dots = \tilde{\theta}_{iJ}$ to all principals when her type is in fact $\theta_i \neq \theta'_i$. Both notions of IC ensure that such consistent lies are not profitable. In addition, UIC imposes IC with respect to *all* inconsistent lies as well; it is the appropriate notion of IC in a one-shot model because agents cannot be punished even if a lie is detected at the end of the game. CIC imposes incentive compatibility over only those messages of i that are in $B_i(\pi)$. CIC is appropriate in the repeated setting because inconsistent messages $(\tilde{\theta}_{i1}, \dots, \tilde{\theta}_{iJ})$ outside $B_i(\pi)$ will induce ‘unexpected’ action profiles with positive probability; when agents are patient, the delayed punishment for this deters such inconsistent messages.¹⁰ Given the notion of IC captured in $K \in \{U, C\}$, a stage-SCF f is said to be induced by $\pi \in \Pi^K$ if $f(\theta) = \pi(\theta)$ for all $\theta \in \Theta$.

4.2.2 Lower bound on a player’s equilibrium payoff

In order to characterize all feasible equilibria in the repeated game $(G, \Gamma)^\infty(\delta)$, SCFs relative to Γ , it is important to lower bound each player’s equilibrium payoff. As discussed earlier, we consider the continuation equilibria where the profile of direct mechanisms induced by agents’ messages from principals’ mechanisms satisfies CIC; hence we denote by w_j^C the principal j ’s minmax value in this repeated game.

¹⁰The reader might wonder if we could further weaken incentive compatibility by using statistical tests to check if the agents reported distribution of types closely approximates the actual distribution. This is addressed later.

We start by introducing some notation. Let

$$\Pi_{-j}^C(\alpha_j) := \{\pi_{-j} \in \Pi_{-j} \mid (\alpha_j, \pi_{-j}) \text{ satisfies } CIC\}.$$

Let Ψ_{-j}^j be the set of all mappings φ_{-j}^j that assigns to each action in $\mathcal{A}_j$ a profile of DMs for principals except j so that $\varphi_{-j}^j(\alpha_j) \in \Pi_{-j}^C(\alpha_j)$ for all $\alpha_j \in \mathcal{A}_j$. Let $u_j^C(\gamma_{-j}, \gamma_j)$ be the lowest stage payoff that principal j receives in a continuation equilibrium at (γ_{-j}, γ_j) . Since principals can change their mechanisms following a principal's deviation, principal j 's minmax value is

$$w_j^C := \min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\gamma_j \in \Gamma_j} u_j^C(\gamma_{-j}, \gamma_j). \quad (10)$$

Most importantly, in the repeated game, we can express this lower bound in terms of actions and DMs as follows.

Theorem 3 *For every principal j , the minmax value w_j^C satisfies*

$$w_j^C = \max_{\alpha_j \in \mathcal{A}_j} \min_{\pi_{-j} \in \Pi_{-j}^C(\alpha_j)} \mathbb{E}_\mu[u_j(\pi_{-j}(\theta), \alpha_j, \theta)]. \quad (11)$$

Off the path following j 's deviation, the following properties are satisfied:

1. w_j^C is achieved if each principal ℓ ($\neq j$) offers an action-reporting DM (ADM) where agents are asked to report their types and principal j 's action;
2. principal j cannot do any better by offering a mechanism in Γ_j than he does by offering an action in $\mathcal{A}_j$.

Proof. See Appendix. ■

This is one of our key results. In a sense it is the main result, because the folk theorem stands on it. It can be decomposed into two separate results, one about mechanisms and the other about minmax payoffs. The first notes that nothing much more complicated than a direct mechanism is needed to mete out punishments. Note that there is no *a priori* restriction on mechanisms; this is an equilibrium result. The second result is that the calculation of the minmax value reduces to the calculation of a maxmin value of a simpler game with a much more restricted space of mechanisms, where it is just a linear programming problem.

The first step in proving this lemma is to use the agents to restrict a deviating principal's set of best responses; this is where the agents' patience plays a critical role. Once we are able to restrict the deviating principal to a simple action, the other principals just need to extract this information from the agents and to offer the worst possible profile of direct mechanisms for the deviating principal.

Of course, off the path following principal j 's deviation, the other principals can offer recommendation mechanisms instead of ADMs. However, reporting the deviating principal's action to a non-deviating principal is simpler than recommending a DM that each non-deviating principal should implement. In a one-shot game, non-deviating principals cannot force agents to induce always the same action from a deviating principal j 's mechanism γ_j regardless of their types. The fact that agents can be punished in the future in the repeated game makes it possible to give incentives to them to induce the same action $g_j(\gamma_j)$ from γ_j if j deviated. This implies that a deviating principal can be punished more severely. As shown later, this, together with a weaker notion of IC (CIC) in the repeated game, makes j 's minmax value in the repeated game with Γ lower than that in the one-shot game with Γ .

Agents' minmax values should be clearly identified because they need to be punished when induced from j 's mechanism γ_j is not $g_j(\gamma_j)$. It is much simpler to identify agent i 's minmax value w_i^C than a principal's minmax value. Let $u_i^C(\gamma)$ be the lowest stage payoff that agent i receives in a continuation equilibrium at γ . Then, w_i^C is defined as $w_i^C := \min_{\gamma \in \Gamma} u_i^C(\gamma)$. Agents' communication with principals in a continuation equilibrium induces a profile of CIC DMs. Therefore, we have

$$w_i^C = \min_{\pi \in \Pi^C} \mathbb{E}_\mu[u_i(\pi(\theta), \theta)]. \quad (12)$$

4.2.3 A folk theorem for competing mechanisms

Now we are ready to establish the folk theorem with i.i.d. private types. What stage-SCFs and payoff profiles can we support in a perfect Bayesian equilibrium (PBE) of $(G, \Gamma)^\infty(\delta)$? The stage-SCF $f \in \mathcal{F}$ is strictly individually rational (SIR) (w.r.t. $\mu \in \Delta(\Theta)$) if each player ℓ gets an expected payoff above w_ℓ^C :

$$\mathbb{E}_\mu[u_\ell(f(\theta), \theta)] > w_\ell^C \text{ for all } \ell \in \mathcal{I} \cup \mathcal{J}. \quad (13)$$

Define

$$\mathcal{F}^C(\mu) := \{f^* \in \Delta(\mathcal{F}) \mid f \in \text{supp}(f^*) \text{ is SIR w.r.t. } \mu \text{ and is induced by } \pi \in \Pi^U\}$$

We make the standard full dimensionality assumption **(FD)** that set of expected payoffs is full dimensional, i.e. $\dim[u(\mathcal{F}^C(\mu))] = J + I$.

The theorem below shows that any (correlated) SCF $f^* \in \mathcal{F}^C(\mu)$ is supportable in a PBE of $(G, \Gamma)^\infty(\delta)$, provided players are sufficiently patient. Given the realization of the public correlation device, suppose that $f \in \mathcal{F}$ in the support of f^* is the SCF that needs to be supported in equilibrium. Principals offer a profile of DMs $\pi = (\pi_1, \dots, \pi_J)$ such that $\pi_j = f_j$ for all j . If principals continue offering π , play a truthful continuation equilibrium in which agents report their true types. If principal j unilaterally deviates at time t , it is observed at the end of the period; at $t + 1$ the other principals offers the ADMs $\bar{\tau}_{-j}^j$ with the majority rule in (24) to punish j . Off the path following principal j 's deviation, principal j can best respond with an action if agents follow the above protocol.

Theorem 4 (Folk Theorem) *Consider i.i.d. types with distribution $\mu \in \Delta\Theta$. Under the standard full dimensionality assumption on the expected payoffs in a model with interdependent values,*

1. *any (correlated) SCF $f^* \in \mathcal{F}^C(\mu)$ is the outcome of a PBE of $(G, \Gamma)^\infty(\delta)$ relative to Γ for high δ ;*
2. *DMs suffice on path, while off the path following a deviation by principal j , all other principals employ ADMs while j offers a simple action, i.e. constant mechanism.*

Proof. See the appendix. ■

Principals' ability to commit to more complex mechanisms (e.g. ADMs) than DMs off the path prevents deviations from DMs; every principal simply relies on DMs on the path, correctly foreseeing all strategic interactions based on ADMs off the path. Therefore, in an equilibrium of the repeated game, we should observe only DMs.

4.3 Comparison: one-shot vs. repeated games

For comparison between one-shot and repeated games, let us first describe the set of SCFs that are supported by equilibria in the one-shot game. Even when an agent's inconsistent type messages are detected, she cannot be punished in the one-shot game. Therefore, the proper notion of IC is UIC. Recall that $u_j^1(\gamma_{-j}, \gamma_j)$ in (8), i.e., the lower bound for principal j 's equilibrium payoff, is the payoff for principal j in the worst continuation equilibrium at $\gamma = (\gamma_{-j}, \gamma_j)$. Let $\Pi^1(\gamma)$ denote the set of all profiles of UIC DMs that can be induced by all continuation equilibria of the one-shot game at $\gamma \in \Gamma$, the superscript 1 referring to the “one-shot” game. Then, in this one-shot game, the lower bound for principal j 's equilibrium payoff w_j^1 specified as

$$w_j^1 = \min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\gamma_j \in \Gamma_j} u_j^1(\gamma_{-j}, \gamma_j), \text{ where } u_j^1(\gamma) := \min_{\pi \in \Pi^1(\gamma)} \mathbb{E}_\mu [u_j(\pi(\theta), \theta)]. \quad (14)$$

To formulate the set of SCFs that are supported by equilibria of the one-shot game, we say that an SCF f is SIR for principals (w.r.t. $\mu \in \Delta\Theta$) if

$$\mathbb{E}_\mu [u_j(f(\theta), \theta)] > w_j^1 \quad \forall j \in \mathcal{J};$$

Define

$$\mathcal{F}^1(\mu) := \{f^* \in \Delta(\mathcal{F}) : f \in \text{supp}(f^*) \text{ is SIR } \forall j \in \mathcal{J} \text{ and is induced by } \pi \in \Pi^U\}. \quad (15)$$

The set of (correlated) SCFs supported in equilibria of the one-shot game (G, Γ) is the closure of $\mathcal{F}^1(\mu)$, denoted by $\bar{\mathcal{F}}^1(\mu)$. This does not provides a characterization of equilibrium allocations in the one-shot game because w_j^1 is not given in terms of model primitives.

The lower bound of agent i 's equilibrium payoff is not well defined in the one-shot game in terms of model primitives. Because it is the lowest possible equilibrium payoff in the one-shot game, it can be specified as

$$w_i^1 = \min_{\pi \in \bar{\mathcal{F}}^1(\mu)} \mathbb{E}_\mu [u_i(\pi(\theta), \theta)]. \quad (16)$$

Give the lower bound of principal j 's payoff w_j^C that can be supported in a PBE

of the repeated game, we can compare it with that in a PBE of the one-shot game. In contrast to the case with no private information on agents' types, the lower-bound of players' equilibrium payoff in the repeated game is generally lower than that in the one-shot game when there is private information on agents' types.

Theorem 5 *For any principal j , $w_j^C \leq w_j^1$ and, for any agent i , $w_i^C \leq w_i^1$.*

Proof. First, consider principal j 's lower bound. From (22), we know that

$$w_j^C = \min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\alpha_j \in \mathcal{A}_j} u_j^C(\gamma_{-j}, \alpha_j),$$

where $u_j^C(\gamma_{-j}, \alpha_j)$ is the lowest payoff that principal j receives in a continuation equilibrium at (γ_{-j}, α_j) in the repeated game.

Because the proper notion of IC is UIC in the one-shot game, $\Pi_{-j}^1(\gamma_j, \alpha_{-j})$ denote the set of all profiles of UIC DMs for principals except j that can be induced by a continuation equilibrium of the one-shot game at (γ_{-j}, α_j) . Following the notation in (14), let us define $u_j^1(\gamma_{-j}, \alpha_j)$ as

$$u_j^1(\gamma_{-j}, \alpha_j) := \min_{\pi_{-j} \in \Pi_{-j}^1(\gamma_j, \alpha_{-j})} \mathbb{E}_\mu [u_j(\pi_{-j}(\theta), \alpha_j, \theta)], \quad (17)$$

that is, principal j 's lowest payoff in a continuation equilibrium at (γ_{-j}, α_j) in the one-shot game. Because CIC is a weaker notion than UIC, any profile of UIC DMs in π_{-j} can be supported in a continuation equilibrium at (γ_{-j}, α_j) in the repeated game. It implies that $u_j^C(\gamma_{-j}, \alpha_j) \leq u_j^1(\gamma_{-j}, \alpha_j)$ and hence

$$w_j^C = \min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\alpha_j \in \mathcal{A}_j} u_j^C(\gamma_{-j}, \alpha_j) \leq \min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\alpha_j \in \mathcal{A}_j} u_j^1(\gamma_{-j}, \alpha_j) \quad (18)$$

Because any single action is strategically equivalent to a constant mechanism in Γ_j that assigns the same action regardless of agents' messages, we have that

$$\min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\alpha_j \in \mathcal{A}_j} u_j^1(\gamma_{-j}, \alpha_j) \leq \min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\gamma_j \in -j} u_j^1(\gamma_{-j}, \gamma_j) = w_j^1. \quad (19)$$

Combining (18) and (19) yields $w_j^C \leq w_j^1$ for all j .

Now consider agent i 's lower bound. Note that $\bar{\mathcal{F}}^1(\mu) \subset \Pi^U$ because UIC is the notion of IC in the one-shot game. Because $\Pi^U \subset \Pi^C$, we then have $\bar{\mathcal{F}}^1(\mu) \subset \Pi^C$.

Therefore, we can conclude that

$$w_i^C = \min_{\pi \in \Pi^C} \mathbb{E}_\mu[u_i(\pi(\theta), \theta)] \leq \min_{\pi \in \bar{\mathcal{F}}^1(\mu)} \mathbb{E}_\mu[u_i(\pi(\theta), \theta)] = w_i^1.$$

■

Clearly, all SCFs that can be supported in a PBE of the one-shot game must be supported in a PBE of the repeated game. Therefore, the set of all SCFs that can be supported in a PBE of the repeated game is

$$\mathcal{F}^*(\mu) := \mathcal{F}^C(\mu) \cup \bar{\mathcal{F}}^1(\mu).$$

Therefore, it is larger than $\bar{\mathcal{F}}^1(\mu)$.

Theorem 6 *Any SCF f^* in $\mathcal{F}^*(\mu)$ can be supported in a PBE of $(G, \Gamma)^\infty(\delta)$ for high δ as follows.*

1. *For any SCF f^* in $\mathcal{F}^C(\mu)$, we can apply Theorem 4 (Folk Theorem) for the proof.*
2. *Any $f^* \in \mathcal{F}^*(\mu) \setminus \mathcal{F}^C(\mu)$ can be supported as a PBE when every principal offers a recommendation mechanism proposed by Yamashita.*

Proof. Any SCF in $\mathcal{F}^C(\mu)$ is shown to be supported via the Folk theorem argument (Theorem 4) with DMs on the path and ADMs off the path: In this case, each pplayer ℓ 's equilibrium payoff is strictly greater than w_ℓ^C .

How about payoffs in $\mathcal{F}^*(\mu) - \mathcal{F}^C(\mu)$? This is the case where $w_\ell^C = w_\ell^1$ for some $\ell \in \mathcal{I} \cup \mathcal{J}$, then we cannot apply Theorem 4. The reason is that the player specific punishments constructed in Lemma 2 for the Folk theorem do not work in this setting because if a player is already at his/her static lower bound w_ℓ^1 and it is the same as his dynamic lower bound w_ℓ^C , we cannot lower it further for punishment. In this case, non-deviating principals offer Yamashita's recommendation mechanisms every period. This is a repetition of Yamashita's static game. ■

Note that the lower bound in $\mathcal{F}^*(\mu)$ is still w_ℓ^C for all $\ell \in \mathcal{I} \cup \mathcal{J}$. Except for the knife-edge case of $w_\ell^C = w_\ell^1$ for some ℓ , we can apply the Folk theorem to show how to support any $f^* \in \mathcal{F}^*(\mu)$ in a PBE of the repeated game by using *simple*

mechanisms such as DMs on the path (and ADMs off the path). This is because principals do not need to stipulate how to punish a competing principal's deviation in their current mechanisms.

Notably, in the repeated game, principal j cannot do any better than offering a simple action off the path following his deviation. Together with it, relaxing the notion of IC lowers the lower bound of principals' equilibrium payoffs. (i.e., $w_j^C \leq w_j^1$). Relaxing the notion of IC also lowers the lower bound of agents' equilibrium payoffs (i.e., $w_i^C \leq w_i^1$). Furthermore, a weaker notion of IC admits additional profiles of DMs on the equilibrium path. Therefore, repeated interaction between players makes it possible to support more social choice functions.

5 Discussion

Our paper investigates the repeated game when principals can commit themselves to one-shot or short-term mechanisms, while the on-going relation among principals and agents governs the long-term relationship without explicit contractual obligations. We now discuss which of our assumptions was made for the analytical simplicity and can be relaxed.

5.1 Markov types

Our results in the case with incomplete information are based on the assumption that types are i.i.d. draws. A more general assumption would be that the types form a Markov chain with the transition matrix P and an initial distribution $\mu^0 \in \Delta\Theta$. We conjecture that theorem 4 can be extended to the case where types evolve according to independent irreducible aperiodic Markov chains.

5.2 Observability

We assume that, at the end of each period, mechanisms and actions are observable to both principals and agents; so DMs are enough on the equilibrium path. However our results extend without any technical changes if only agents, but not principals, observe all actions and mechanisms at the end of each period. Since principals do not observe mechanisms and actions, on-path mechanisms must ask agents to report not

only their own types but also the identity of any principal who may have deviated; the corresponding message space is simply the product of the type space with the set $\{0, 1, \dots, J\}$, where 0 signifies that no principal has deviated; principals base their action on the report sent by a majority. This creates no additional difficulty for our folk theorem because there is common knowledge of who the deviator is.¹¹ However, the principal's action choice in his equilibrium mechanism depends on only agents' type reports but not the identity of the deviating principal that they report. Although messages on the deviator's identity is purely cheap talk, no agent can affect play by unilaterally deviating from reporting a deviator.

5.3 Two agents

A principal's deviation is observable because mechanisms and actions are observable. However, an agent's deviation may not be easily detected when there are only two agents. Off the path following principal j 's deviation, non-deviators offer ADMs. If two agents' reports on j 's action to principal ℓ are not consistent, principal ℓ knows that at least one agent has deviated. However, other principals do not because agents' action reports to principal ℓ are not observable by them. This makes it difficult to punish agents' deviation because all principals should punish agents together after their deviation. Of course, if there are three or more agents, we can always assume agents' truthful reporting of j 's action to other principals.

Therefore, if there is a tie on agents' reports on j 's action to principal ℓ off the path following j 's deviation, principal ℓ should let other principals know about that so that they can all punish agents in the following periods. Principal ℓ may choose an action that is not chosen in any circumstances given agents' truthful reports. For any given $(\tilde{\alpha}_\ell^j, \tilde{\theta}_\ell)$, if $\#\{i : \tilde{\alpha}_{i\ell}^j = \alpha_j\} = \#\{i : \tilde{\alpha}_{i\ell}^j = \alpha'_j\} = I/2$, then

$$\bar{\tau}_\ell^j(\tilde{\alpha}_\ell^j, \tilde{\theta}_\ell) = \hat{\alpha}_\ell \notin \{\alpha_\ell = \bar{\varphi}_\ell^j(\alpha_j)(\theta) \mid \forall (\alpha_j, \theta) \in \mathcal{A}_j \times \Theta\} \quad (20)$$

Suppose that there are two agents and they both report principal j 's true action. If an agent deviates to report a false action, then their action messages are inconsistent. Suppose that one reports α_j and the other α'_j . Principal ℓ assigns action $\hat{\alpha}_\ell$ regardless

¹¹Off the path following a principal's deviation, non-deviating principals offer an extended DM that ask agents about their types, the deviating principal's action, and for the identity of any player who might have deviated from carrying out the on-going punishment.

of agents' type messages. The formulation in (20) implies that $\hat{a}_\ell$ is never assigned by principal j if both agents report principal j 's true action. It implies that by taking $\hat{a}_\ell$, principal ℓ informs the other principals and agents that inconsistent messages on principal j 's action were reported. Then, agents can be punished subsequently. Therefore, if the action space is sufficiently rich to ensure the existence of such an action $\hat{a}_\ell$ that satisfied (20), then agents' truthful reporting on principal j 's action can be enforceable even with only two agents or in the case of a tie.

6 Conclusion

Our paper studies a very natural game—repeated period-by-period contracting among multiple principals and agents. The one-shot version of this model is well studied. While it is indeed a very natural model, its use has been hampered by its complexity at two levels. First, one might want to compute the equilibrium allocation, say by computing lower bounds on the payoffs. Second, leaving aside a characterisation, partial or full, we might simply be interested in understanding how to support a particular equilibrium allocation. Unfortunately neither has been tackled in the one-shot game, largely because all manner of deviations need to be deterred contemporaneously. Usually repeating a game introduces complexity: in the world of competing mechanisms, the model becomes more tractable as we move from the one-shot game to the repeated game with patient players—pointing this out is our main contribution. The key insight is that agents can be used to play the role of monitor, and they in turn can be monitored by the principals. Thanks to the participation of agents in the punishment of deviating principals, it is possible to lower the minmax value. This lowered value is achieved by agents reducing the complex mechanisms of a deviating principal to simple actions. Furthermore, computing the minmax values reduces to a simple programming problem.

References

- [1] Abreu, D., D. Dutta and L. Smith (1991), “The Folk Theorem for Repeated Games: A NEU Condition,” *Econometrica*, 62(4), 939-948.

- [2] Abreu, D. and H. Matsushima (1992), “Virtual Implementation in Iteratively Undominated Strategies I: Complete Information,” *Econometrica*, 60, 993-1008.
- [3] Bergemann, D. and S. Morris (2005), “Robust Mechanism Design,” *Econometrica*, 73(6), 1771–1813.
- [4] Breiman, L. (1991), Probability, Society for Industrial and Applied Mathematics, 1991.
- [5] Crámer, J. and R. McLean (1988), “Full Extraction of the Surplus in Bayesian and Dominant Strategy Auctions,” *Econometrica*, 56, 1247–1257.
- [6] Epstein L. and M. Peters (1999), “A Revelation Principle for Competing Mechanisms,” *Journal of Economic Theory* 88(1), 119-160.
- [7] Ghosh, S. and S. Han (2013), “Repeated Competition among Principals,” mimeo, McMaster University.
- [7] McAfee, P. (1993), “Mechanism Design by Competing Sellers,” *Econometrica*, 61(6), 1281-1312.
- [8] Myerson, R. (1979), “Incentive Compatibility and the Bargaining Problem,” *Econometrica*, 47(1), 61-7.
- [9] Peck, J. (1997), “A Note on Competing Mechanisms and the Revelation Principle,” Ohio State University, mimeo.
- [10] Peters, M., (2014), “Competing Mechanisms,” *Canadian Journal of Economics*, 47(2), 373-397
- [11] Peters, M., and C. Troncoso-Valverde, (2013), “A Folk Theorem for Competing Mechanisms,” *Journal of Economic Theory*, 148(3), 953-973.
- [12] Szentes, B. (2010), “A note on mechanism games with multiple principals and three or more agents by T. Yamashita,” Discussion paper. London School of Economics, London, UK
- [13] Yamashita, T. (2010), “Mechanism games with multiple principals and three or more agents,” *Econometrica*, 78(2), 791-801.

A Appendix

A.1 Proof of Theorem 3

Proof. The set of mechanisms Γ_j is sufficiently large to include all possible constant mechanisms. Because a constant mechanism is equivalent to an action, it is clear that

$$\min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\gamma_j \in \Gamma_j} u_j^C(\gamma_{-j}, \gamma_j) \geq \min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\alpha_j \in \mathcal{A}_j} u_j^C(\gamma_{-j}, \alpha_j) \quad (21)$$

We now show that (21) holds with equality. The image set of a mechanism γ_j is denoted by $\text{im}(\gamma_j)$. Take an selection $g_j : \Gamma_j \rightarrow \mathcal{A}_j$ where $g_j(\gamma_j) \in \text{im}(\gamma_j)$ is an arbitrary action. Then, for every possible γ_j that principal j can offer off the path following his deviation, let agents send messages that induce $g_j(\gamma_j) \in \text{im}(\gamma_j)$ in the continuation equilibrium irrespective of their own types. . Therefore, this type of communication behavior completely neutralizes principal j 's ability to make his action choice contingent on agents' messages: Principal j cannot do any better by offering a complex mechanism in Γ_j than by offering a single action in $\mathcal{A}_j$:

$$u_j^C(\gamma) = u_j^C(\gamma_{-j}, g_j(\gamma_j)) \quad \forall \gamma.$$

This implies that

$$\min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\gamma_j \in \Gamma_j} u_j^C(\gamma_{-j}, \gamma_j) = \min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\gamma_j \in \Gamma_j} u_j^C(\gamma_{-j}, g_j(\gamma_j)) \leq \min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\alpha_j \in \mathcal{A}_j} u_j^C(\gamma_{-j}, \alpha_j).$$

Combining this with (21) we find that inequality (21) holds with equality:

$$\min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\gamma_j \in \Gamma_j} u_j^C(\gamma_{-j}, \gamma_j) = \min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\alpha_j \in \mathcal{A}_j} u_j^C(\gamma_{-j}, \alpha_j) \quad (22)$$

Now we can focus on principal j 's deviation to a single action in $\mathcal{A}_j$ without loss of generality. For any given single action α_j that principal j takes, for the other principals, punishing principal j in a continuation equilibrium is equivalent to choosing a profile of their CIC DMs conditional on α_j . Therefore, principal j 's lowest possible payoff conditional on α_j can be realized if other principals can implement $\bar{\varphi}_{-j}^j(\alpha_j) = \{\bar{\varphi}_\ell^j(\alpha_j)\}_{\ell \neq j} \in \Pi_{-j}^C(\alpha_j)$, where $\bar{\varphi}_\ell^j(\alpha_j)$ is principal ℓ 's DM and $\bar{\varphi}_{-j}^j(\alpha_j)$ is

defined as

$$\bar{\varphi}_{-j}^j(\alpha_j) \in \arg \min_{\pi_{-j} \in \Pi_{-j}^C(\alpha_j)} \mathbb{E}_\mu[u_j(\pi_{-j}(\theta), \alpha_j, \theta)]. \quad (23)$$

The DM $\bar{\varphi}_{-j}^j(\alpha_j)$ can be implemented if principal ℓ offers an ADM $\bar{\tau}_\ell^j$ in which he commits to take DM $\bar{\varphi}_{-j}^j(\alpha_j)$ when a majority of agents report α_j in addition to their types. Let us define a profile of ADMs for principals except j that can implement $\bar{\varphi}_{-j}^j$. Let $\bar{\tau}_\ell^j$ be principal ℓ 's ADM where each agent i 's message space is $T_{i\ell} = \mathcal{A}_j \times \Theta_i$. Let $T_\ell = \times_{i \in \mathcal{I}} T_{i\ell}$. Then, $\bar{\tau}_\ell^j$ is a mapping from T_ℓ into $\mathcal{A}_\ell$. Let $(\tilde{\alpha}_{i\ell}^j, \tilde{\theta}_{i\ell})$ denote a pair of messages that agent i sends to principal ℓ , where $\tilde{\alpha}_{i\ell}^j$ is principal j 's action that agent i reports to principal ℓ and $\tilde{\theta}_{i\ell}$ is the type that she reports to him as her type. In this ADM, principal ℓ asks agents about principal j 's action and their types. Let $\tilde{\alpha}_\ell^j$ be a profile of all agents' messages on principal j 's action and $\tilde{\theta}_\ell$ a profile of all agents' messages on their types. Then, $\bar{\tau}_\ell^j(\tilde{\alpha}_\ell^j, \tilde{\theta}_\ell) \in \mathcal{A}_\ell$ denotes principal ℓ 's action when $(\tilde{\alpha}_\ell^j, \tilde{\theta}_\ell)$ is a profile of agents' messages. We use the notation $\bar{\tau}_\ell^j(\tilde{\alpha}_\ell^j, \cdot)$ to denote a DM that is assigned when a profile of agents' messages on principal ℓ 's action. For the implementation of $\bar{\varphi}_{-j}^j$, we use an ADM with the *majority rule*:

$$\bar{\tau}_\ell^j(\tilde{\alpha}_\ell^j, \cdot) = \bar{\varphi}_{-j}^j(\alpha_j) \text{ if } \#\{i : \tilde{\alpha}_{i\ell}^j = \alpha_j\} > I/2 \quad (24)$$

Assuming Γ_{-j} is sufficiently large to include ADMs, we have the following equality given agents' truthful reporting to principals, except j , who offer ADMs with the majority rule in (24):

$$\min_{\gamma_{-j} \in \Gamma_{-j}} \max_{\alpha_j \in \mathcal{A}_j} u_j^C(\gamma_{-j}, \alpha_j) = \max_{\alpha_j \in \mathcal{A}_j} \mathbb{E}_\mu[u_j(\bar{\varphi}_{-j}^j(\alpha_j)(\theta), \alpha_j, \theta)]. \quad (25)$$

Applying the definition of $\bar{\varphi}_{-j}^j(\alpha_j)$ in (23), it is straightforward to see that

$$\max_{\alpha_j \in \mathcal{A}_j} \mathbb{E}_\mu[u_j(\bar{\varphi}_{-j}^j(\alpha_j)(\theta), \alpha_j, \theta)] = \max_{\alpha_j \in \mathcal{A}_j} \min_{\pi_{-j} \in \Pi_{-j}^C(\alpha_j)} \mathbb{E}_\mu[u_j(\pi_{-j}(\theta), \alpha_j, \theta)]. \quad (26)$$

(22), (23), and (26) lead to (11).

Finally, we need to show what action $g_j(\gamma_j)$ in the image of γ_j leads to the lowest possible payoff for principal j in the worst continuation equilibrium off the path where principal j 's mechanism is γ_j . Because the other principals can implement CIC DMs $\bar{\varphi}_{-j}^j(\alpha_j)$ that punishes principal j most severely given his action α_j ,

the action $g_j(\gamma_j)$ that agents induce from principal j 's mechanism γ_j in the worst continuation equilibrium is

$$g_j(\gamma_j) \in \min_{\alpha_j \in \text{im}(\gamma_j)} \mathbb{E}_\mu[u_j(\bar{\varphi}_{-j}^j(\alpha_j)(\theta), \alpha_j, \theta)]. \quad (27)$$

Therefore, agents send messages that induce an action $g_j(\gamma_j)$ defined in (27) when principal j offers γ_j off the path following his deviation. $\blacksquare$

A.2 Proof of Theorem 4

Caveat: In discussions related to folk theorems, generic players are denoted by i and j unless explicitly noted otherwise.

Definition 2 Fix a SIR SCF $f \in \mathcal{F}$ that is induced by $\pi \in \Pi^C$. A family of vectors $\{\beta^1, \dots, \beta^{J+I}\} \subset S^C$ is said to be a PSP (Player-Specific Punishment) for the target payoff $v = u(f)$ if it satisfies the following properties $\forall i, j \in \mathcal{I} \cup \mathcal{J}$:

1. *strict individual rationality (SIR):* $\beta_j^i > w_j$;
2. *target payoff domination:* $\beta_j^i < v_j$;
3. *payoff asymmetry (PA):* $\beta_i^i < \beta_i^j$ if $i \neq j$.

Lemma 2 Fix a SIR SCF $f \in \mathcal{F}$ that is induced by $\pi \in \Pi^C$. There exists a family of $I + J$ profiles of mechanisms $\{\pi^i : \Theta \rightarrow \mathcal{A} \mid i \in \mathcal{I} \cup \mathcal{J}\}$ such that each π_k^i is one-to-one and the family $\{\beta^i := \mathbb{E}_\mu u(\pi^i(\theta)) \mid i \in \mathcal{I} \cup \mathcal{J}\}$ is a PSP for $v = u(f)$.

Proof. Given full-dimensionality we can construct¹² a PSP $\{\bar{\beta}^i \mid i \in \mathcal{I} \cup \mathcal{J}\}$. Since $\bar{\beta}^i \in u(\mathcal{F}^C(\mu))$, by construction there exists a family of DMs $\{\bar{\pi}^i \mid i \in \mathcal{I} \cup \mathcal{J}\}$ such that $\bar{\beta}^i := \mathbb{E}_\mu u(\bar{\pi}^i(\theta))$. Since properties 1, 2 and 3 above rely on strict inequalities it is easy to see that there is an $r > 0$ such that any family $\{y^i \mid y^i \in B(\bar{\beta}^i, r)\}$ is also a PSP. If any such $\bar{\pi}_k^i$, for $i \in \mathcal{I} \cup \mathcal{J}$ and $k \in \mathcal{J}$, is not one-to-one, replace it with a new DM π_k^i as follows. (If any $\bar{\pi}_k^i$ is one-to-one, set $\pi_k^i(\theta) = \bar{\pi}_k^i(\theta)$.) Fix any enumeration of the type-space $\Theta = \{\theta^1, \dots, \theta^L\}$. Let

$$\pi_k^i(\theta^1) := \bar{\pi}_k^i(\theta^1),$$

¹²See Abreu, Dutta and Smith (1991).

and for $l \geq 1$ pick an arbitrary element

$$\pi_k^i(\theta^{l+1}) \in B(\bar{\pi}_k^i(\theta^l), r) \setminus \{\pi_k^i(\theta^1), \dots, \pi_k^i(\theta^l)\}.$$

Now define $\beta^i := \mathbb{E}_\mu[u(\pi^i(\theta))]$ for all $i \in \mathcal{I} \cup \mathcal{J}$. ■

Proof of the Theorem.

Fix any (correlated) SCF $f^* \in \mathcal{F}^C(\mu)$. Suppose that f in the support of f^* is a SCF that needs to be supported given the realization of the public correlation device; let $v := u(f)$ denote the corresponding payoff.

Let $\bar{u} := \max_{\mathcal{I}, \mathcal{A}, \Theta} |u_i(a, \theta)|$ be the least upper bound on the payoffs in the stage-game. Find a parameter $q \in (0, 1)$ such that

$$\bar{u}(1 - q) < \beta_i^i(2 - q) - w_i^C \quad \forall i \in \mathcal{I} \cup \mathcal{J}. \quad (28)$$

Such a q exists because at $q = 1$ this inequality becomes $0 < \beta_i^i - w_i^C$, which is satisfied by construction of β^i (see Lemma 2). Use Lemma 2 to find a family of profiles of direct mechanisms $\{\pi^i : \Theta \rightarrow \mathcal{A} \mid 1 \leq i \leq I + J\}$ such that each π_k^i is one-to-one and the family $\{\beta^i := \mathbb{E}_\mu u(\pi^i(\theta)) \mid i \in \mathcal{N}\}$ is a PSP for $v = u(f)$.

Strategies are defined by the following rules.

1. Play starts in phase I. Each principal j offers the DM $\pi_j^* = f_j$. Agents report the actual type θ_i^t to all principals at time t . If principal j deviates unilaterally (offers a mechanism other than π_j^*), agents play a one-shot continuation equilibrium in the current period; play moves to phase II $_j$ from the next period. If π^* is offered but the action profile chosen in the current period does not belong to $\hat{\mathcal{A}}(\pi^*)$, all players can infer that some agent deviated; if one agent i can be identified as the deviator, move to phase II $_i$; otherwise move to phase II $_j$ with probability $1/I$ for each agent j (not for principals).
2. Phase II $_j$ proceeds as follows for $1 \leq j \leq J$. Let $g^j : \Gamma_j \rightarrow \mathcal{A}_j$ be as in equation (27) and $\bar{\varphi}_k^j(\tilde{\alpha}_j)$ be as defined in (23). Principal j can potentially offer any complex mechanism in this phase. When γ_j is offered by principal j , agents send messages to j to induce the action $g^j(\gamma_j) \in \text{im}(\gamma_j)$ irrespective of their types. Each principal $k \neq j$ offers the ADM $\bar{\tau}_k^j$ that asks the agents

to report the action $g^j(\gamma_j)$ and their true type; if a majority of agents reports $g^j(\gamma_j)$, the DM $\bar{\varphi}_k^j(g^j(\gamma_j))$ is used by principal k to determine his own action.

Phase II_i proceeds as follows for $J+1 \leq i \leq J+I$. Principals offer the profile π^i that attains the minmax value w_i^C of agent i . If any player l deviates while in this phase, start phase II_l ; otherwise, switch to phase III_i with probability $1 - q \in (0, 1)$ after each period spent in phase II_i .

If a deviation is detected, then (i) start phase II_l if player l can be identified as the deviator w.p. 1; (ii) if there is no unique player who can be identified as the deviator, move to phase II_j with probability $1/I$ for each agent j (not for principals). If there is no deviation, switch to phase III_j with probability $1 - q \in (0, 1)$ independently across time after each period spent in phase II_j .

3. In phase III_j , for $1 \leq j \leq J+I$, principal k offers $\bar{\pi}_k^j \in \Pi_k$, which has expected payoff vector β^j ; agents report truthfully. Remain forever in this phase, unless any player l deviates unilaterally and triggers phase II_l .

VERIFICATION OF EQUILIBRIUM:

A deviation by a principal is observable by all players. However not all deviations by agents can be detected with probability 1. If principals offer a profile of DMs π satisfying CIC, no consistent but false type report is profitable; any profitable inconsistent type report by agent i is detected with a minimum probability

$$\min_{(\tilde{\theta}_{ij})_{j \in (\tilde{\Theta})^J}} \mu \left\{ \theta_{-i} \in \Theta_{-i} \mid (\pi_1(\tilde{\theta}_{i1}, \theta_{-i}), \dots, \pi_J(\tilde{\theta}_{iJ}, \theta_{-i})) \notin \hat{\mathcal{A}}(\pi) \right\}.$$

For any profile π , let $p_{\min}(\pi)$ be the minimum of the above probabilities over all agents i . Now define

$$p_{\min} := \min \{ p_{\min}(\pi^*), p_{\min}(\bar{\pi}^1), \dots, p_{\min}(\bar{\pi}^{I+J}). \}$$

By the one-shot deviation principle, it suffices to show that the proposed strategy is unimprovable, i.e. no one-shot deviation by any player i from any phase is profitable.

Let L_j^i denote player j 's expected utility from the beginning of phase II_i . First, with $j = i$, player i 's lifetime (discounted average) payoff in phase II_i is defined

recursively as $L_i^i = (1 - \delta)w_i^C + \delta(qL_i^i + (1 - q)\beta_i^i)$, so that

$$L_i^i = \frac{(1 - \delta)w_i^C + \delta(1 - q)\beta_i^i}{1 - \delta q}. \quad (29)$$

Note that $L_i^i \rightarrow \beta_i^i$ as $\delta \rightarrow 1$. When calculating L_j^i for $j \neq i$ we have to bound the expected utility on both sides:

$$(1 - \delta)(-\bar{u}) + \delta(qL_j^i + (1 - q)\beta_j^i) \leq L_j^i \leq (1 - \delta)\bar{u} + \delta(qL_j^i + (1 - q)\beta_j^i).$$

As $\delta \rightarrow 1$ it is easy to check that $L_j^i \rightarrow \beta_j^i$.

1. Phase III_i for $1 \leq i \leq I + J$: Note that the mechanisms used in this phase are all one-to-one by construction; this means that any deviator, principal or agent, is commonly known to all players, and the quantity $p_{\min}$ plays no role. From the definitions it is clear that the difference in the lifetime payoffs to one-shot deviation and conformity is at least

$$(1 - \delta)\bar{u} + \delta L_i^i - \beta_i^i = (1 - \delta) \left[\bar{u} - \frac{(1 + \delta - \delta q)\beta_i^i - \delta w_i^C}{1 - \delta q} \right] \quad (30)$$

using (29). An immediate implication of inequality (28) defining q is that (30) is strictly negative for all δ close to 1, so that i cannot profitably deviate from Phase III_i. Since $\beta_j^i > \beta_j^j \forall j \neq i$, it is immediate that $(1 - \delta)\bar{u} + \delta L_i^i - \beta_i^i$ for such δ and therefore players $j \neq i$ do not have a profitable one-shot deviation either from Phase III_i.

2. Phase II_i for $1 \leq i \leq I + J$: Any agent j will not deviate from II_i if

$$(1 - \delta)\bar{u} + \delta\{(1 - p)L_j^i + \frac{p}{I} \sum_{k \neq j} L_j^k + \frac{p}{I} L_j^j\} \leq L_j^i,$$

where $p \geq p_{\min}$ is the prob with a deviation is detected when agent j deviates. Hence a sufficient condition is

$$\frac{\max_k \beta_j^k - \min_k \beta_j^k}{\max_k \beta_j^k - \beta_j^j} < \frac{p_{\min}}{I}.$$

Caveat: When we choose the betas we have to make sure that for some $\varepsilon > 0$

we have $\beta_j^j + 2\varepsilon = \beta_j^i = \beta_j^k$ for all $k, i \neq j$ and that the perturbation to make one-to-one is less than ε .

Since deviation by any principal j are obs w.p. 1, it is easier to check that deviations are not profitable.

3. Phase I for $1 \leq i \leq I + J$: A deviation by a principal is identified with probability 1. Note that unlike a usual repeated game this probability may be strictly lower than 1 in case an agent sends an inconsistent message outside her $B_i(\pi^*)$. (Recall that inconsistent messages within $B_i(\pi^*)$ cannot be detected and are deterred by IC.)

$$(1 - \delta)\bar{u} + \delta\{(1 - p)v_j + \frac{p}{I} \sum_{k \neq j} L_j^k + \frac{p}{I} L_j^j\} \leq v_j,$$

which simplifies to

$$\sum_{k \neq j} \beta_j^k + \beta_j^j < Iv_j,$$

which is satisfied because $v_j > \beta_j^i$ (target payoff domination) for all i and j ; strategies in phase I are therefore unimprovable.

In sum, for high δ , the posited strategy profile is unimprovable after all histories, and hence is an equilibrium. ■