

# Continuous implementation with direct revelation mechanisms<sup>☆</sup>

Yi-Chun Chen<sup>a</sup>, Manuel Mueller-Frank<sup>b</sup>, Mallesh M. Pai<sup>c,\*</sup>

<sup>a</sup> *Department of Economics and Risk Management Institute, National University of Singapore, Singapore*

<sup>b</sup> *IESE Business School, University of Navarra, Spain*

<sup>c</sup> *Department of Economics, Rice University, United States of America*

Received 13 January 2020; final version received 24 October 2021; accepted 31 January 2022

Available online 15 February 2022

---

## Abstract

We investigate how a principal's knowledge of agents' higher-order beliefs impacts their ability to robustly implement a given social choice function. We adapt a formulation of Oury and Tercieux (2012): a social choice function is continuously implementable if it is partially implementable for types in an initial model and "nearby" types. We characterize when a social choice function is *truthfully* continuously implementable, i.e., using game forms corresponding to direct revelation mechanisms for the initial model. Our characterization hinges on how our formalization of the notion of nearby preserves agents' higher order beliefs. If nearby types have similar higher order beliefs, truthful continuous implementation is roughly equivalent to requiring that the social choice function is implementable in strict equilibrium in the initial model, a very permissive solution concept. If they do not, then our notion is equivalent to requiring that the social choice function is implementable in unique rationalizable strategies in the initial model. Truthful continuous implementation is thus very demanding without non-trivial knowledge of agents' higher order beliefs.

© 2022 Elsevier Inc. All rights reserved.

---

<sup>☆</sup> The authors thank Rahul Deb, Matt Jackson, Maher Said, Satoru Takahashi, Olivier Tercieux, Siyang Xiong and Muhamet Yildiz for helpful comments. We would especially like to thank the editor, Tilman Börgers, an anonymous Associate editor and three anonymous referees for their comments and suggestions on earlier drafts. Pai was partially supported by NSF CCF-1763349. Chen gratefully acknowledges support from Singapore Ministry of Education Academic Research Fund Tier 1 (R-122-000-296-115).

\* Corresponding author.

E-mail addresses: [ecsycc@nus.edu.sg](mailto:ecsycc@nus.edu.sg) (Y.-C. Chen), [mmuellerfrank@iese.edu](mailto:mmuellerfrank@iese.edu) (M. Mueller-Frank), [mallesh.pai@rice.edu](mailto:mallesh.pai@rice.edu) (M.M. Pai).

JEL classification: D82; D83

Keywords: Continuous implementation; Robust implementation; Contagion; Higher-order beliefs

## 1. Introduction

The literature on Robust Mechanism Design, starting with the seminal work of Bergemann and Morris (2005) studies settings where the designer does not perfectly understand the information structure among agents. It investigates the design of mechanisms that perform robustly well across various information structures among agents that the principal considers possible. In this paper, our aim is to isolate how a desire for robustness impacts a principal who is solely unsure about agents' higher-order beliefs, i.e. beliefs of agents about each other's beliefs etc. Distinguished contributions in the game theory literature inform us that predictions in a given strategic situation can be very sensitive to agents' higher-order beliefs (e.g. Rubinstein (1989) or Weinstein and Yildiz (2007)). Our question thus concerns how these higher-order beliefs play a role when the principal can design the game among the agents.

We start from a standard Bayesian implementation setting: there are finite sets of agents, states and alternatives, and there is a commonly known information structure that describes the information of the agents. The planner would like to (partially) implement a given social choice function, i.e. a function from profiles of types to alternatives. In this case, any Bayesian incentive compatible social choice function can be partially implemented with a direct revelation mechanism. But what if the principal is unsure about the exact information structure among agents, but would nevertheless like the social choice function to be partially implemented “close to” a reference information structure? Formally, we adapt the formulation of Oury and Tercieux (2012) and revisit the question of when a social choice function is continuously implementable.<sup>1</sup>

Our main results characterize when a social choice function is *truthfully* continuously implementable, i.e., using game forms corresponding to direct revelation mechanisms for the initial model. One way to interpret our restriction is that it formalizes conditions under which a principal who believes a baseline information structure and therefore uses a direct revelation mechanism is nevertheless able to implement his desired social choice function when he is “slightly” wrong. Under this interpretation, our notion of truthful continuous implementation is a robustness check to the standard revelation principle—we build on this interpretation by presenting results on the set of continuously implementable social choice functions. An alternate interpretation is that by limiting the message space, we rule out “detail-free” mechanisms that simply elicit these details from the agents and then proceed akin to standard mechanism design. Such mechanisms, it may be argued, obey the letter but not the spirit of a robustness exercise.<sup>2</sup>

Intuitively, the characterization depends on the underlying topology with respect to which we demand continuity. We study two well understood topologies in this setting. The

<sup>1</sup> Our paper substantially builds off their work, we defer a fuller discussion of the details of their work, the closely related characterization of Oury (2015), and other related papers to Section 6, after we have formally stated our own results.

<sup>2</sup> Of course, a principal may opt for a different “simple” mechanism rather than a direct revelation mechanism. To that end, note that while our results are formally stated for direct revelation mechanisms, our proof techniques apply to any mechanism where the equilibrium in the baseline is full-range, i.e. for every message available to any agent, there is some type of agent in the baseline information structure which sends that message. We expand on this observation below after presenting our formal results.

first, the product topology, only preserves lower order beliefs. It is the topology studied in Oury and Tercieux (2012) (also, the topology implicitly used in Rubinstein (1989) and explicitly appealed to in Weinstein and Yildiz (2007)). The second is the uniform-weak topology of Chen et al. (2010), which preserves higher-order beliefs. The latter is studied for two reasons. Firstly, we would argue, this is of independent interest: being a finer topology, continuity with respect to this topology captures a weaker notion of robustness. Conceptually, one can argue that these capture two disparate ways an information structure can be close to a given information structure: the latter involves agreement at all arbitrarily higher-order beliefs, while the former topology only constrains lower-order beliefs. Second, at a more technical level, our results in the latter are a building block for our results in the former—we detail this further below in Section 1.1. In Section 5, we develop an example of a standard government natural resource auction setting to motivate these topologies.

At a high level, our findings can be summarized thus: settings like the latter, where despite not knowing the exact information structure, the principal has information about the agents' higher-order beliefs, are not much more constraining than the baseline of *exact* knowledge of the information structure. By contrast, if the agents' higher-order beliefs may be arbitrary, then the principal is severely restricted.

Further, we show that a “revelation principle” applies for the latter notion. In that setting, if a social choice function can be continuously implemented, it can be truthfully continuously implemented by a direct revelation mechanism. A revelation principle does not obtain in the more general setting. Requiring this stronger notion, therefore, may necessitate the use of more complex mechanisms to continuously implement some social choice functions (in particular, mechanisms containing messages that are not sent by any type in equilibrium in the baseline information structure considered by the principal). Further, we provide a partial characterization of continuous implementation in this setting, and thus explain the gap between continuous implementation and truthful continuous implementation.

### 1.1. Model and results

Let us now describe the setting and our results more formally. There are finite sets of agents, states and alternatives.<sup>3</sup> There is given a social choice function of interest. There is a baseline information structure that the principal considers. The actual information structure that obtains among agents is unknown to the principal. We wish to understand when the social choice function can be truthfully continuously implemented: i.e. in any (epistemic) model that embeds the baseline model, there is an equilibrium of the direct revelation mechanism such that the baseline types report their types truthfully (resulting in the desired social choice function), and further the strategy of closeby types converges. We term this requirement *truthful continuous implementation* (the additional modifier of “truthful” to the notion of Oury and Tercieux (2012) reflecting our restriction to the truthful equilibrium of a direct revelation mechanism).

We study continuity with respect to two topologies on types. The first, the product topology, places no restrictions on agents' higher-order beliefs. We show that under this topology, truthful continuous implementation is equivalent to requiring that the social choice function be implementable with a mechanism such that, in the baseline model, each agent has a unique rationalizable action, and the desired alternative of the social choice function obtains if each agent

<sup>3</sup> Throughout, we assume a richness condition on the environment: see Section 2.3 for details.

plays this unique rationalizable action (Theorem 1). The second, the uniform-weak topology, (see e.g. Monderer and Samet (1989) and Chen et al. (2010)) is roughly a topology that preserves higher-order beliefs. We show that under this topology, a social choice function is truthfully continuously implementable if and only if it can be implemented in Strict equilibrium in the baseline model (Theorem 2).

Finally, we shed some light on the gap between continuous implementation and truthful continuous implementation. We show that a social choice function is continuously implementable with respect to the uniform-weak topology if and only if it is truthfully continuously implementable with respect to the uniform-weak topology (Theorem 3). Therefore a revelation principle holds for continuous implementation with respect to the uniform-weak topology. However, we show that one does not get a revelation principle with respect to the product topology.<sup>4</sup> In particular, our methods show something stronger—if a social choice function is not truthfully continuously implementable, but is continuously implementable, then the implementing mechanism must necessarily have messages that are not being sent at the baseline.<sup>5</sup>

At a technical level, we would like to highlight our characterization results in the product topology. To get some intuition for this result, recall the work of Weinstein and Yildiz (2007). They consider a *given* game of incomplete information. They assume a form of richness: for each player, and each action of that player, there exists a “crazy type” whose preferences make that action strictly dominant. Their main result is to show that for any action  $a$  that is rationalizable for a (normal) type in the game, there exist close-by types in the product topology for whom that action is the unique rationalizable action. The possibility of aforementioned crazy types is used to start a contagion process, with the strict dominance used to break ties. In an implementation setting, this assumption of crazy types is not well grounded, since the game form is chosen by the planner and therefore not fixed a priori. Further, we are after a partial equilibrium result, i.e. there exists one equilibrium of the game with the desired properties.<sup>6</sup>

Instead our result in the product topology builds off of our result in the uniform-weak topology. Closeness in the uniform-weak topology implies closeness in the product topology. By our results in the former, we know that the social choice function must be implementable in Strict Bayes-Nash Equilibrium. Recall further that we are considering implementation with DRMs, i.e. for every message an agent could send there is a corresponding type: in other words, the equilibrium has full range. Strict equilibrium implies that for that type it is a strict best response for him to send the corresponding message. We use these types as a substitute for the crazy types described above—these are sufficient since we are indeed arguing the existence (or lack thereof) of a single equilibrium.

Take any rationalizable strategy  $s_i$  for a player  $i$ . We construct a sequence of types that converge to the baseline type in the product topology for which this strategy is the unique best

<sup>4</sup> We can give a partial characterization of continuous implementation with respect to the product topology: we show that any continuously implementable social choice function must be strictly rationalizable implementable. The converse need not be true.

<sup>5</sup> Formally, we show a setting with two possible types in the baseline for each agent where the direct-revelation mechanism (which, by definition, has two messages per agent) is not continuously implementable (i.e., the desired social choice function is not truthfully continuously implementable). However, we show by construction an indirect mechanism with *three* messages for each agent which does continuously implement the desired social choice function. In particular, each of the two baseline types for each agent has a unique rationalizable action in this mechanism, so for each agent, there is a message that is not sent in equilibrium by any baseline type.

<sup>6</sup> In this sense, there is a tighter connection between our results and those of Weinstein and Yildiz (2004), we discuss the details after we introduce our formal results. See also Weinstein and Yildiz (2011).

response, in a manner similar to Weinstein and Yildiz (2007) (and also Weinstein and Yildiz (2004): see discussion after the proof of the theorem). Roughly, put most of the mass of  $i$ 's beliefs on the fact the others will play the strategies that rationalize  $s_i$ , and a small probability of the type corresponding to the strategy  $s_i$ . The latter makes this a strict best response. Therefore, at *any* Bayes-Nash Equilibrium of the incomplete information game in this model, these constructed types must be playing the rationalizable strategy  $s_i$ . From the fact that the social choice function is continuously implementable, therefore, we have rationalizable implementation as desired.

The paper is organized as follows. Section 2 defines the model. Section 3 characterizes truthful continuous implementation. Section 4 studies the original continuous implementation of Oury and Tercieux (2012) in this setting and the gap between the two. Section 5 develops an application in the context of natural resource auctions and explains the implications of our results. Section 6 discusses the related literature and connections.

## 2. Model

There is a state of the world  $\theta \in \Theta$ , unknown to the planner. There is a set of alternatives  $A$ . Unless otherwise stated, both  $A$  and  $\Theta$  are finite. There is a finite set of  $I$  agents. Agent  $i$  has a utility function  $u_i : A \times \Theta \rightarrow \mathbb{R}$ . Sometimes, we might refer directly to the implied ordinal preferences over alternatives, with the standard notations  $\succ_{i,\theta}$  for the strict part of the preference of agent  $i$  at state  $\theta$ ,  $\sim_{i,\theta}$  for indifferences, and  $\succeq_{i,\theta}$  for weak preference.

### 2.1. Epistemic preliminaries

A model  $\mathcal{T}$  is a pair  $(T, \kappa)$  where  $T = T_1 \times T_2 \times \cdots \times T_I$  is a countable type space and  $\kappa_{t_i} \in \Delta(\Theta \times T_{-i})$  denotes the associated beliefs for each  $t_i \in T_i$ .

Given a type  $t_i$  in a model  $(T, \kappa)$ , we can compute the first-order belief of  $t_i$  (i.e., his belief about  $\Theta$ ) by setting  $t_i^1$  equal to the marginal distribution of  $\kappa_{t_i}$  on  $\Theta$ . We can also compute the second-order belief of  $t_i$  (i.e., his belief about  $(\theta, t^1)$ ) by setting

$$t_i^2[E] = \kappa_{t_i} \left[ \left\{ (\theta, t_{-i}) : (\theta, t_i^1, t_{-i}^1) \in E \right\} \right], \forall E \subset \Theta \times (\Delta(\Theta))^I.$$

We can compute the entire hierarchy of beliefs  $(t_i^1, t_i^2, \dots, t_i^k, \dots)$  iteratively.

Now, write  $X^0 = \Theta$  and for each  $k \geq 1$ :  $X^k = [\Delta(X^{k-1})]^I \times X^{k-1}$ . Observe that  $t_i^k \in \Delta(X^{k-1})$  for every  $k \geq 1$ . Let  $d^0$  be the discrete metric on  $\Theta$  and  $d^1$  be the Prohorov distance on 1st-order beliefs  $(\Delta(\Theta))^I$ .<sup>7</sup> Then, recursively, for any  $k \geq 2$ , endow  $\Delta(X^{k-1})$  with the Prohorov distance  $d^k$  where  $X^{k-1}$  is endowed with the sup-metric induced by  $d^0, d^1, \dots, d^{k-1}$ . Mertens and Zamir (1985) construct the universal type space  $T_i^* \subset \times_{k=0}^\infty \Delta(X^k)$ . The universal type space has the property that  $t_i = (t_i^1, t_i^2, \dots) \in T_i^*$  if there exists some type  $t'_i$  in some model such that  $t_i$  and  $t'_i$  have the same  $n$ -th-order belief for every  $n$ . Endowed with the product topology,  $T_i^*$  is a compact metrizable space and admits a homeomorphism  $\kappa_i^* : T_i^* \rightarrow \Delta(\Theta \times T_{-i}^*)$ .

<sup>7</sup> For a metric space  $(X, \rho)$ , the Prohorov distance between any two  $\mu, \mu' \in \Delta(X)$  is

$$\inf\{\gamma > 0 : \mu'(E) \leq \mu(E^\gamma) + \gamma \text{ for every Borel set } E \subseteq X\},$$

$$\text{where } E^\gamma = \{x \in X : \inf_{y \in E} \rho(x, y) < \gamma\}.$$

We say that a sequence of types  $\{t_{i,n}\}_{n=1}^{\infty}$  converges uniform-weakly to a type  $t_i$  if:

$$d_i^{\text{uw}}(t_{i,n}, t_i) \equiv \sup_{k \geq 1} d_i^k(t_{i,n}^k, t_i^k) \rightarrow 0.$$

Moreover, write  $d^{\text{uw}}(t_n, t) \rightarrow 0$  if  $d_i^{\text{uw}}(t_{i,n}, t_i) \rightarrow 0$  for each  $i$ .<sup>8</sup> Similarly, a sequence of types  $\{t_{i,n}\}_{n=1}^{\infty}$  converges in the product topology to a type  $t_i$  if

$$d_i^{\text{p}}(t_{i,n}, t_i) \equiv \sum_{k=1}^{\infty} 2^{-k} d_i^k(t_{i,n}^k, t_i^k) \rightarrow 0.$$

Again, write  $d^{\text{p}}(t_n, t) \rightarrow 0$  if  $d_i^{\text{p}}(t_{i,n}, t_i) \rightarrow 0$  for each  $i$ .

Following Oury and Tercieux (2012), for two models  $\mathcal{T} = (T, \kappa)$  and  $\mathcal{T}' = (T', \kappa')$ , we will write  $\mathcal{T} \supset \mathcal{T}'$  if  $T \supset T'$ , and for  $t_i \in T'_i : \kappa_{t_i}[E] = \kappa'_{t_i}[(\Theta \times T'_{-i}) \cap E]$  for any measurable  $E \subset \Theta \times T_{-i}$ .

The principal considers a baseline model which we denote by  $\overline{\mathcal{T}} = (\overline{T}, \overline{\kappa})$ . We assume that the baseline model is finite, i.e.,  $|\overline{T}| < \infty$ ; and nonredundant, i.e., no distinct types  $t_i$  and  $t'_i$  in  $\overline{T}_i$  induce the same hierarchy of beliefs. For instance, this includes as a special case the standard mechanism design setting with a common prior over payoff-relevant types. More precisely, we may set  $\Theta = \times_{i \in I} \Theta_i$ ,  $T_i = \Theta_i$ , and each  $\kappa_{t_i}$  is induced from a common prior  $\mu \in \Delta(\Theta)$  such that  $\text{marg}_{\Theta_i} \mu[\theta_i] > 0$  for each  $\theta_i$ , i.e.,  $\kappa_{t_i}[(\theta_i, \theta_{-i}, t_{-i})] = 1_{\{\theta_i = t_i, \theta_{-i} = t_{-i}\}} \mu(\theta_{-i} | \theta_i)$ .

## 2.2. Mechanisms and notion of implementation

There is a subset of types  $\overline{T}_0 \subseteq \overline{T}$  that the principal cares about. A social choice function (SCF) is a mapping  $f : \overline{T}_0 \rightarrow A$ . In general  $\overline{T}_0 = \overline{T}$ , but in some examples we may have strict containment. Assume also that  $\{t_i\} \times \text{supp} \kappa_{t_i} \subset \overline{T}_0$  for every  $t_i \in \overline{T}_i$  (the reason for this support condition is so that social choice function is well-defined for every profile that every type considers possible).

A mechanism, denoted  $\mathcal{M} = (M, g)$  is a message space  $M_i$  for each player  $i$ , with  $M \equiv \times_i M_i$ , and an outcome function  $g : M \rightarrow A$ . A countable (respectively, finite) mechanism is one where the message space  $M$  is countable (respectively, finite) in cardinality. Given a mechanism  $\mathcal{M}$  and a model  $\mathcal{T}$ , we write  $U(\mathcal{M}, T)$  for the induced incomplete information game. A Bayes-Nash Equilibrium (BNE) is a strategy profile  $(\sigma_i)_{i \in I}$  with  $\sigma_i : T_i \rightarrow \Delta(M_i)$  such that for  $t_i \in T_i$ , each message  $m_i \in \text{supp } \sigma_i(t_i)$  maximizes the expected payoff of agent  $i$  with respect to the opponents' strategy profile  $\sigma_{-i}$ .

A direct revelation mechanism is defined as is standard, i.e. the message space of every player equals the set of types the principal considers possible in the baseline model, and the outcome function is denoted  $g : \times_i \overline{T}_i \rightarrow A$ . We can now define truthful continuous implementation in this setting:

**Definition 1.** We say  $f$  is *truthfully continuously implementable w.r.t. metric  $d$*  if there is a direct revelation mechanism  $g$  such that for any model  $\mathcal{T} \supset \overline{\mathcal{T}}$ , there is a (possibly mixed) BNE  $\sigma$  in the game  $U(g, \mathcal{T})$  such that for every  $t \in \overline{T}_0$ :

<sup>8</sup> See Chen et al. (2010) for further details about this topology.

- a.  $g(t) = f(t)$ , and,
- b. for any sequence  $\{t_n\} \subset T$  with  $d(t_n, t) \rightarrow 0$ ,  $\sigma(t_n) \rightarrow \delta_t$ .<sup>9</sup>

Definition 1 is directly comparable to the definition of continuous implementation of Oury and Tercieux (2012) (Definition 2 in their paper)—see Definition 2 below for their definition in our notation. Note that truthful continuous implementation is more demanding than continuous implementation in two ways. Firstly, it fixes the form of the mechanism used: the former restricts attention to direct revelation mechanisms where the latter considers general mechanisms. Secondly, it demands robustness of a specific equilibrium of this mechanism (i.e., the truth-telling equilibrium), whereas the latter focuses on outcomes.<sup>10</sup>

**Definition 2.** Given any SCF  $f$ , mechanism  $\mathcal{M} = (M, g)$ , and model  $\mathcal{T} = (T, \kappa)$  with  $\mathcal{T} \supset \overline{\mathcal{T}}$ , say that a mixed-strategy Bayes-Nash Equilibrium (BNE)  $\sigma$  continuously implements (resp. strictly continuously implements)  $f$  in  $U(\mathcal{M}, \mathcal{T})$  w.r.t. a metric  $d$  if

- a.  $\sigma|_{\overline{\mathcal{T}}}$  is a pure-strategy Bayes-Nash Equilibrium (resp. strict Bayes-Nash equilibrium) in  $U(\mathcal{M}, \overline{\mathcal{T}})$ ;
- b. For any  $t \in \overline{T}_0$ ,  $g(\sigma(t_n)) \rightarrow f(t)$  for any sequence of type profiles  $\{t_n\} \subset T$  with  $d(t_n, t) \rightarrow 0$ .

We say that  $f$  is *continuously implementable* (resp. *strictly continuously implementable*) w.r.t. metric  $d$  if there is a mechanism  $\mathcal{M} = (M, g)$  such that for any model  $\mathcal{T} \supset \overline{\mathcal{T}}$ , there is an equilibrium which continuously implements (resp. strictly continuously implements)  $f$  w.r.t.  $d$  in  $U(\mathcal{M}, \mathcal{T})$ .

### 2.3. Reduced normal forms and a richness assumption

A recurring issue in our setting is breaking indifferences, since we have no transfers. To get results within a classical implementation setting we therefore need a richness assumption.<sup>11</sup> In order to introduce our assumption, first consider the following standard definition of strategic equivalence adapted to our setting.

**Definition 3.** For a DRM  $g$ , we say  $t_i$  is *strategically equivalent* to  $t'_i$  for an agent  $i$  if agent  $i$  is indifferent between the two reports regardless of the state and others' reports, i.e.:

$$\forall t_{-i}, \theta : g(t_i, t_{-i}) \sim_{i, \theta} g(t'_i, t_{-i}).$$

In light of this we can define the reduced normal-form of a DRM, again, in line with standard terminology.

<sup>9</sup> Note that the space of messages is finite, and so convergence is in the standard e.g. Euclidean topology in the finite dimensional simplex.

<sup>10</sup> We discuss this further in Section 6.

<sup>11</sup> Our assumption serves the same technical role as the assumption of costly messages in Oury and Tercieux (2012) and local payoff uncertainty in Oury (2015). We discuss those assumptions when we compare to the related literature in Section 6. At a high-level though, the difference is conceptual—our assumption is one that can be verified in the context of the baseline model considered by the principal. Their assumptions are richness assumptions on the elaborations of their model with respect to which continuous implementation is desired.



**Definition 4.** A *reduced normal-form* of a DRM  $g$ , denoted  $\tilde{g}$ , is a mechanism in which all the strategically equivalent messages are identified. For each  $t_i$ , let  $\tilde{t}_i$  denote the message in  $\tilde{g}$  corresponding to the set of messages strategically equivalent to  $t_i$  in  $g$ .

It is possible in the original mechanism  $g$  that two messages are strategically equivalent for some agent  $i$  but deliver different outcomes at some profile of messages from other agents, i.e. the mechanism  $\tilde{g}$  is not well defined. The following assumption rules this out.

**Assumption 1.** We say that a DRM  $g$  *admits a reduced normal-form* if  $\tilde{g}$  is well defined, i.e., for an agent  $i$  and any two messages  $t_i$  and  $t'_i$  which are strategically equivalent,  $g(t_i, \cdot) = g(t'_i, \cdot)$ .

This is reminiscent of the non-bossiness assumption of Satterthwaite and Sonnenschein (1981), which is often invoked in social choice/allocation settings. Roughly, it requires that if an agent changing his report (all else equal) changes the selected alternative, then the agent cannot be indifferent between the two alternatives. However, non-bossiness is standardly defined only for private-value settings, so we do not expound further.

This assumption is novel and therefore perhaps not well understood. Observe that the following simple richness assumption implies that Assumption 1 is always satisfied: in particular this assumption is purely on the environment rather than Assumption 1 which is on the environment and the desired social choice function  $f$ .

**Assumption 2.** For every agent  $i$  and any two alternatives  $a, a' \in A$ , there is some  $\theta$  such that agent  $i$  is not indifferent between  $a$  and  $a'$  under  $\theta$ .

This latter assumption may not be appropriate for some settings of interest. For example, in a private-good allocation setting, agents may be always indifferent between alternatives that only differ in the allocations of other agents. Even here, however, the desired social choice function  $f$  may be such that Assumption 1 is satisfied, even though the environment does not satisfy Assumption 2 (note that the latter assumption does not depend on the social choice function  $f$ ).

To see this consider the following private-good, private-value allocation setting. There are three agents 1, 2, 3, and three alternatives 1, 2, 3, with each alternative to be thought of as the corresponding agent getting the good. Each agent  $i$  has a type  $t_i \in [0, 1]$  which is their value for receiving the good, and an outside option of 0 for not receiving the good, with  $\theta = (t_1, t_2, t_3)$ ,  $\Theta = [0, 1] \times [0, 1] \times [0, 1]$ . Observe first that in this setting, Assumption 2 is not satisfied—e.g. agent 1 is always indifferent between alternatives 2 and 3. However, note that the social choice function which assigns the good efficiently,  $f(t_1, t_2, t_3) = \arg \max_i (t_1, t_2, t_3)$  is such that any DRM  $g$  that implements it must satisfy Assumption 1—an agent's report will sometimes affect her own allocation. In fact, in this example, there are no strategically equivalent messages.

In what follows, we invoke the weaker Assumption 1. The reader may mentally substitute the stronger Assumption 2 if they prefer. Either way, we emphasize that either of these assumptions are directly verifiable on the primitives of the model.

### 3. Characterizing truthful continuous implementation

Our main result in this section is a characterization of the set of truthfully continuously implementable social choice functions in the product topology. The following definition of interim correlated rationalizable messages (cf. Dekel et al. (2007)) will be useful:



**Definition 5.** Let  $R_i^\infty(t_i, \mathcal{M})$  denote the set of *interim correlated rationalizable* messages of type  $t_i$  in  $\mathcal{M}$  defined as follows:

Let  $R_i^0(t_i, \mathcal{M}) = M_i$ . Inductively, for each  $k \geq 1$ , a message  $m_i \in R_i^k(t_i, \mathcal{M})$  iff there is some  $\mu \in \Delta(\Theta \times T_{-i} \times M_{-i})$  such that

$$\mathbf{R1:} \quad m_i \in \arg \max_{m'_i} \int_{\Theta \times M_{-i}} u_i(m'_i, m_{-i}, \theta) \operatorname{marg}_{\Theta \times M_{-i}} \mu[d\theta, m_{-i}];$$

$$\mathbf{R2:} \quad \operatorname{marg}_{\Theta \times T_{-i}} \mu = \kappa_{t_i};$$

$$\mathbf{R3:} \quad \mu \left( \left\{ (\theta, t_{-i}, m_{-i}) : m_{-i} \in R_{-i}^{k-1}(t_{-i}, \mathcal{M}) \right\} \right) = 1.$$

$$\text{Then, } R_i^\infty(t_i, \mathcal{M}) \equiv \bigcap_{k=1}^\infty R_i^k(t_i, \mathcal{M}).$$

We can now define implementation in unique rationalizable action profile:

**Definition 6.** Let  $g$  be a DRM that admits a reduced normal-form. We say  $f$  is implementable in the unique rationalizable action profile in the reduced normal-form  $\tilde{g}$  if for every  $t \in \bar{T}$ ,  $R^\infty(t, \tilde{g}) = \{\tilde{t}\}$ .

Note that this definition is slightly stronger than rationalizable implementation: the latter only requires that every rationalizable action profile results in the desired alternative, while in addition, we require that the implementing mechanism have a unique rationalizable strategy for each type.

**Theorem 1.** Suppose that Assumption 1 holds. An SCF  $f$  is truthfully continuously implementable w.r.t.  $d^p$  by a DRM  $g$  if and only if it is implementable in unique rationalizable action profile in  $\tilde{g}$ .

Since this proof is fairly involved, a high level overview may be useful to help orient the reader. Sufficiency is fairly straightforward—if  $\tilde{g}$  implements  $f$  in unique rationalizable action, then  $g$  truthfully continuously implements  $f$ —this follows straightforwardly from the upper hemicontinuity of the rationalizable correspondence.

The nontrivial direction is therefore necessity, i.e. to show that if an SCF  $f$  is truthfully continuously implementable (in the product topology) then  $f$  must be implementable in the unique rationalizable action in the sense of Definition 6.

As a key building block we use our characterization of truthful continuous implementation in uniform-weak topology below (Theorem 2). Combined with Corollary 1 this tells us that an SCF  $f$  is truthfully continuously implementable w.r.t. the uniform-weak topology if and only if it is implementable in Strict Bayes-Nash Equilibrium in the “reduced normal form.”<sup>12</sup> From this fact, and the fact that the uniform-weak topology is finer than the product topology, we have that if  $f$  is truthfully continuously implementable (in the product topology) then  $f$  is implementable in Strict Bayes-Nash Equilibrium.

Recall further that we are considering implementation with DRMs, i.e. for every message an agent could send there is a corresponding type: in other words, the equilibrium has full range. Strictness implies that for the type corresponding to a particular message it is a strict best re-

<sup>12</sup> Note that the latter is well-defined by Assumption 1. As we discussed earlier, one could have invoked the stronger, but easier to verify Assumption 2.

sponse for him to send the corresponding message. We use this fact as a substitute for the costly messages of Oury and Tercieux (2012) or the local payoff uncertainty of Oury (2015).

Take any type  $t'_i$  which is a rationalizable report for a player  $i$  of type  $t_i \in \bar{T}_i$ . We can construct a sequence of types  $t''_i$  that converge to  $t_i$  in the product topology for which reporting  $t'_i$  is the unique best response, in a manner similar to Weinstein and Yildiz (2007) (and also Weinstein and Yildiz (2004)). Roughly, put most of the mass of  $i$ 's beliefs on the fact the others will play the strategies that rationalize  $t_i$ , and a small probability that the type is  $t'_i$ . The latter makes reporting  $t'_i$  a strict best response. Therefore, at *any* Bayes-Nash Equilibrium of the incomplete information game in this model, these constructed types must be playing the rationalizable message  $t'_i$ . Since the social choice function is continuously implementable, therefore, we have rationalizable implementation as desired.

### 3.1. Uniform-weak topology

We now introduce our characterization of truthful continuous implementation in the uniform-weak topology. As we pointed out above, this is useful as a stepping stone to the characterization in the product topology. Since continuity with respect to the uniform-weak topology captures a weaker notion of robustness, these results may be of independent interest. To state and prove our characterization, we introduce two more terms. We say that DRM  $g$  *strictly rewards truth-telling* at type  $t_i$  over type  $t'_i$  for agent  $i$  if

$$\sum_{(\theta, t_{-i}) \in \Theta \times \bar{T}_{-i}} [u_i(g(t_i, t_{-i}), \theta) - u_i(g(t'_i, t_{-i}), \theta)] \bar{\kappa}_{t_i}[(\theta, t_{-i})] > 0.$$

We say that (reporting)  $t_i$  always weakly dominates  $t'_i$  for agent  $i$  in DRM  $g$  if

$$\forall (\theta, t_{-i}) \in \Theta \times \bar{T}_{-i} : u_i(g(t_i, t_{-i}), \theta) - u_i(g(t'_i, t_{-i}), \theta) \geq 0.$$

The following lemma is key to our characterization.

**Lemma 1.** *If an SCF  $f$  is truthfully continuously implementable by a DRM  $g$  with respect to  $d^{uw}$ , then, for every agent  $i$  and any pair of agent  $i$ 's types  $t_i$  and  $t'_i$ , either  $g$  strictly rewards truth-telling at  $t_i$  over  $t'_i$ ; or  $t_i$  always weakly dominates  $t'_i$  in  $g$ .*

Suppose there exists an agent  $i$  and a pair of types  $t_i$  and  $t'_i$  such that  $g$  neither strictly rewards truth-telling at  $t_i$ , nor does  $t_i$  always weakly dominate  $t'_i$ . This in particular means that there is some state  $\theta'$  and some profile of other agents' reports  $t'_{-i}$  at which agent  $i$  strictly prefers to report  $t'_i$  over  $t_i$ . We show that there exists a sequence of perturbations which converges to  $t_i$  in the uniform-weak topology, such that each type in this sequence uniquely prefers to report  $t'_i$  in the DRM. Roughly speaking, these are types that put a small mass on the state that the type is  $\theta'$  and the other agents' types are  $t'_{-i}$ , but are otherwise identical to  $t_i$ . Thus the conditions described in Lemma 1 are a necessary condition for truthful continuous implementation in this setting.

Our main characterization of truthful continuous implementation follows:

**Theorem 2.** *An SCF  $f$  is truthfully continuously implementable by a DRM  $g$  with respect to  $d^{uw}$  if and only if  $g(t) = f(t)$  for all  $t \in \bar{T}_0$  and, for every agent  $i$  and any pair  $t_i$  and  $t'_i$ , either  $g$  strictly rewards truth-telling at  $t_i$  over  $t'_i$ ; or  $t_i$  is strategically equivalent to  $t'_i$  for agent  $i$ .*

The proof of this theorem is easy to describe. The necessity of our condition is straightforward in light of Lemma 1. If  $g$  does not strictly reward truth-telling at  $t_i$  over  $t'_i$ , then by the condition of the Lemma,  $t_i$  must always weakly dominate  $t'_i$ . But then  $g$  cannot strictly reward truth-telling at  $t'_i$  over  $t_i$  either. This must imply that  $t'_i$  also always weakly dominates  $t_i$ , which implies that  $t_i$  and  $t'_i$  are strategically equivalent. We show the sufficiency of our condition constructively.

**Corollary 1.** *Suppose that Assumption 1 holds.  $f$  is truthfully continuously implementable in  $d^{uw}$  if and only if the reduced normal-form DRM  $\tilde{g}$  implements  $f$  in truthful strict BNE in  $U(\mathcal{M}, \overline{T})$ , i.e. if truth-telling is a strict Bayes-Nash equilibrium in the game  $U(\mathcal{M}, \overline{T})$ .*

As an aside we should note that similar permissive results would be achieved if we considered closeness in the strategic topology of Dekel et al. (2006). This follows from a result of Chen et al. (2010) who show that the two topologies are equivalent around finite types (recall that by assumption the baseline model was finite).

#### 4. A revelation principle for continuous implementation?

So far, we have only studied truthful continuous implementation. We now recall the definition of continuous implementation in this setting and consider the relation between continuous implementation and truthful continuous implementation for both topologies.

We begin with a positive result, i.e. that if requiring continuous implementation with respect to the uniform-weak topology, we have a revelation principle.

To state and prove our characterization of continuous implementation, we adapt two definitions to this environment. Fix a mechanism  $\mathcal{M} = (M, g)$ . For agent  $i$ 's type  $t_i$  in  $\overline{T}_i$  and message  $m'_i \in M_i$ , we say that  $f$  strictly rewards  $\overline{\sigma}_i(t_i)$  over  $m'_i$  in a (pure-strategy) BNE  $\overline{\sigma}$  in  $U(\mathcal{M}, \overline{T})$  if

$$\sum_{(\theta, t_{-i}) \in \Theta \times \overline{T}_{-i}} [u_i(g(\overline{\sigma}_i(t_i), \overline{\sigma}_{-i}(t_{-i})), \theta) - u_i(g(m'_i, \overline{\sigma}_{-i}(t_{-i})), \theta)] \overline{\kappa}_{t_i}[(\theta, t_{-i})] > 0.$$

We say that  $\overline{\sigma}_i(t_i)$  always weakly dominates  $m'_i$  in a (pure-strategy) BNE  $\overline{\sigma}$  in  $U(\mathcal{M}, \overline{T})$  if

$$\forall (\theta, t_{-i}) \in \Theta \times \overline{T}_{-i}: u_i(g(\overline{\sigma}_i(t_i), \overline{\sigma}_{-i}(t_{-i})), \theta) - u_i(g(m'_i, \overline{\sigma}_{-i}(t_{-i})), \theta) \geq 0.$$

The following lemma is again the key to our characterization of continuous implementation. The proof is analogous to the proof of Lemma 1.

**Lemma 2.** *If  $\overline{T}_0 = \overline{T}$  and  $f$  is continuously implementable w.r.t.  $d^{uw}$  by mechanism  $\mathcal{M} = (M, g)$ , then there is a pure-strategy BNE  $\overline{\sigma}$  in  $U(\mathcal{M}, \overline{T})$  such that for each agent  $i$ , each type  $t_i$  in  $\overline{T}_i$  and message  $m'_i \in M_i$ , either  $f$  strictly rewards  $\overline{\sigma}_i(t_i)$  over  $m'_i$ ; or  $\overline{\sigma}_i(t_i)$  always weakly dominates  $m'_i$  in BNE  $\overline{\sigma}$ .*

Lemma 2 immediately implies the following characterization (as well as revelation principle) for continuous implementation in  $d^{uw}$ . Denote by  $\tilde{f}$  the reduced normal form of the DRM  $f$ .

**Theorem 3.** *Suppose that Assumption 1 holds and  $\overline{T}_0 = \overline{T}$ .  $f$  is continuously implementable in  $d^{uw}$  if and only if the reduced normal-form DRM  $\tilde{f}$  implements  $f$  in truthful strict BNE in  $U(\mathcal{M}, \overline{T})$ .*

The basic idea of Theorem 3 is analogous to the proof of Theorem 2. The main difference is that we need Assumption 1 to ensure that the reduced normal-form is well-defined. We can then apply similar arguments. Comparing to Theorem 2, we therefore have that, with respect to the uniform-weak topology, a social choice function is continuously implementable iff it is truthfully continuously implementable.

**Proof.** ( $\Rightarrow$ ) Let  $\mathcal{M} = (M, g)$  be a mechanism such that BNE  $\sigma$  continuously implements  $f$ . Consider the direct revelation mechanism  $\mathcal{M}' = (g', \bar{T})$  defined as  $g'(t) = g(\sigma(t))$  for all  $t \in \bar{T}$ . By Lemma 2 and Theorem 2 such a mechanism clearly truthfully continuously implements  $f$ . The implication now follows from Corollary 1.

( $\Leftarrow$ ) By Corollary 1, an SCF  $f$  satisfying this condition is truthfully continuously implementable and therefore trivially, also continuously implementable.  $\square$

#### 4.1. Product topology

In this section, we first show by counterexample that a revelation principle does not apply to continuous implementation with respect to the product topology. In particular, we show an example below in which the direct revelation mechanism does not continuously implement the desired social choice function (in particular, since it is easily verified that this fails the characterization of Theorem 1). We then constructively show that there is a mechanism which contains additional messages and continuously implements the desired social choice function. The example is essentially due to Oury and Tercieux (2012) (working paper version).

There are 2 agents. The set of outcomes is

$$A = \{(x, p_1, p_2) : x \in \{0, 1, 2, 3\}, p_1, p_2 \in \{0, 35, 40\}\}.$$

If  $x = 0$ , the object is not given to either agent,  $x = 1$  or 2 connotes that it was given to the respective agent, while  $x = 3$  implies that neither agent gets the object and both are punished. The  $p_i$ 's correspond to payments from the agents to the principal. Utility functions are quasilinear and the object has a monetary value to each agent.

Each agent is either of type  $\theta_1$  or  $\theta_2$ . Agent  $i$  with type  $\theta_i$  has value 50 for the object, and agent  $i$  with type  $\theta_j$  with  $j \neq i$  has value 30 for the object. Finally, the agents' utility of  $x = 0$  is zero and the punishment outcome  $x = 3$  is equivalent to a value of  $-30$  to agent  $i$  if they are of  $\theta_i$ , and  $-50$  if they are of type  $\theta_j$  for  $j \neq i$ .

The baseline type space of each agent is  $\{\theta_1, \theta_2\}$  with a common prior  $P(\theta_i, \theta_i) = \frac{1-\varepsilon}{2}$  and  $P(\theta_i, \theta_j) = \frac{\varepsilon}{2}$  for  $i = 1, 2$  and  $j \neq i$ . Hence, type  $\theta_1$  believes the other agent is of type  $\theta_1$  with probability  $(1 - \varepsilon)$ , and type  $\theta_2$  with probability  $\varepsilon$ , and  $\theta_2$ 's beliefs are defined analogously. Let  $\bar{T}_i = \{\theta_1, \theta_2\}$  and  $\bar{T}_0 = \bar{T}_1 \times \bar{T}_2$ . That is, the baseline type space has a product form as we assume in Theorem 3. The social choice function that the principal would like to continuously implement is  $f(\theta_i, \theta_i) = (i, 0, 0)$ ,  $f(\theta_i, \theta_j) = (0, 40, 40)$ . Finally pick  $\varepsilon$  small enough that such that  $5 \times (1 - \varepsilon) - 100 \times \varepsilon > 0$ . For the two mechanisms that we present below, 5 is the minimal payoff difference and 100 the maximal payoff difference for distinct outcomes.

**Claim 1.** *This social choice function is not truthfully continuously implementable with respect to  $d^P$ .*

**Proof.** A direct revelation mechanism in this setting has exactly two messages for each agent, one corresponding to each type; moreover, the outcomes are given by  $f$ :

|            |             |             |
|------------|-------------|-------------|
|            | $\theta_2$  | $\theta_1$  |
| $\theta_1$ | (0, 40, 40) | (1, 0, 0)   |
| $\theta_2$ | (2, 0, 0)   | (0, 40, 40) |

Observe that for  $i = 1, 2$ , the agents reporting  $(\theta_i, \theta_i)$  regardless of types is a strict Bayes Nash equilibrium. Hence, both messages are rationalizable for both types. Thus,  $f$  is not implementable in unique rationalizable action profile in the direct revelation mechanism (in which no strategies are strategically equivalent). Therefore the claim follows directly from the characterization of Theorem 1.  $\square$

**Claim 2.** *There exists an indirect mechanism that continuously implements  $f$  with respect to  $d^P$ .*

**Proof.** Consider an indirect mechanism where each agent has 3 possible messages, (Mine, His, Mine+). The outcome is given by the matrix below:

|       |             |             |             |
|-------|-------------|-------------|-------------|
|       | Mine        | His         | Mine+       |
| Mine  | (0, 40, 40) | (1, 0, 0)   | (2, 40, 35) |
| His   | (2, 0, 0)   | (0, 40, 40) | (0, 35, 0)  |
| Mine+ | (1, 35, 40) | (0, 0, 35)  | (3, 0, 0)   |

First, action “His” is strictly dominated by “Mine+” for agent 1 with type  $\theta_1$ . Consequently, “Mine” and Mine+ are strictly dominated by “His” for agent 2 with type  $\theta_1$ . Finally, in the third round, “Mine” is strictly better than “Mine+” for agent 1 with type  $\theta_1$ . Analogous reasoning follows for type  $\theta_2$ . Hence “Mine” is the unique rationalizable action for agent  $i$  with type  $\theta_i$ , and “His” for agent  $i$  with type  $\theta_j$ . Playing this rationalizable action results in the desired social choice function being implemented.

Therefore, the mechanism described above continuously implements the social choice function  $f$  w.r.t.  $d^P$  because the interim correlated rationalizable correspondence is upper-hemi-continuous (see proof of sufficiency of Theorem 1).  $\square$

#### 4.2. A partial characterization for indirect mechanisms

Finally, we provide some results about continuous implementation with respect to the product topology in indirect mechanisms. We assume that  $\overline{T}$  has full support, i.e., for each  $t_i \in \overline{T}_i$ , we have  $\text{supp} \overline{\kappa}_{t_i} = \overline{T}_{-i}$ . Some new definitions are now necessary. We say that  $m_i$  is strategically equivalent to  $m'_i$  for agent  $i$  in BNE  $\overline{\sigma}$  in  $U(\mathcal{M}, \overline{T})$  if

$$\forall (\theta, t_{-i}) \in \Theta \times \overline{T}_{-i} : u_i(g(m_i, \overline{\sigma}_{-i}(t_{-i})), \theta) = u_i(g(m'_i, \overline{\sigma}_{-i}(t_{-i})), \theta).$$

The following assumption is essentially Assumption 1 adapted to indirect mechanisms.<sup>13</sup>

**Assumption 3.** For any agent  $i$  and any two messages  $m_i$  and  $m'_i$  which are strategically equivalent for some BNE  $\overline{\sigma}$  in  $U(\mathcal{M}, \overline{T})$ , we have  $g(m_i, \cdot) = g(m'_i, \cdot)$ .

<sup>13</sup> We note that this assumption is somewhat more opaque—for example there is no natural analog to Assumption 2 which can be easily verified without reference to the social choice function/implementing mechanism. To that end, we should note that we do not consider Assumption 3 as “natural” or “desireable”—it is simply the assumption under which we are able to make some progress understanding continuous implementation with indirect mechanisms.

**Theorem 4.** Suppose that Assumption 3 holds for mechanism  $\mathcal{M}$ . Then,  $f$  is continuously implementable in  $d^p$  if and only if it is strictly continuously implementable in  $d^p$ .

It is worth connecting our results to Oury and Tercieux (2012). Their Theorem 3 shows that any social choice function that is strictly continuously implementable must satisfy a form of monotonicity (formally, strict interim rationalizable monotonicity, see Definition 8 of that paper). The present theorem effectively shows that under Assumption 3, the same implication extends to all continuously implementable social choice functions.

**Definition 7.** Let  $\mathcal{T} = (T, \kappa)$  be a model. Denote by  $W_i^\infty(t_i, \mathcal{M})$  the set of (interim correlated) strictly rationalizable messages of type  $t_i$  in  $U(\mathcal{M}, \mathcal{T})$  defined as follows:

Let  $W_i^0(t_i, \mathcal{M}) = M_i$ . Inductively, for each  $k \geq 1$ , a message  $m_i \in W_i^k(t_i, \mathcal{M})$  iff there is some  $\mu_{-i} \in \Delta(\Theta \times T_{-i} \times M_{-i})$  such that

$$\mathbf{R1}: \{m_i\} = \arg \max_{m'_i} \sum_{\theta, m_{-i}} u_i(m'_i, m_{-i}, \theta) \arg_{\Theta \times M_{-i}} \mu[\theta, m_{-i}];$$

$$\mathbf{R2}: \arg_{\Theta \times T_{-i}} \mu_{-i} = \kappa_{t_i};$$

$$\mathbf{R3}: \mu_{-i} \left( \left\{ (\theta, t_{-i}, m_{-i}) : m_{-i} \in W_{-i}^{k-1}(t_{-i}, \mathcal{M}) \right\} \right) = 1.$$

$$\text{Then, } W_i^\infty(t_i, \mathcal{M}) \equiv \bigcap_{k=1}^\infty W_i^k(t_i, \mathcal{M}).$$

We can now define implementation in strictly rationalizable action profiles:

**Definition 8.** We say  $f$  is implementable in strictly rationalizable action profiles by mechanism  $\mathcal{M}$  if for every  $t \in \bar{T}$ , we have  $g(m) = f(t)$  for every  $m \in W^\infty(t, \mathcal{M})$ .

**Theorem 5.** Suppose that Assumption 3 holds. An SCF  $f$  is continuously implementable w.r.t.  $d^p$  by a finite mechanism only if  $f$  is implementable in strictly rationalizable action profiles by a finite mechanism.

As we pointed out earlier, our proof techniques in Theorem 1 apply to any mechanism such that at the baseline information structure there is an equilibrium which is both full-range and implements our desired social choice function. The desideratum of “equilibrium continuous implementation” would be defined with respect to this equilibrium, by analogy to Definition 1. It should be clear that our characterization of Theorem 1 continues to hold in such a case. The gap between Theorem 1 and Theorem 5 is that the latter allows for mechanisms that contain messages not sent by any type in the baseline information structure (as in the construction of Claim 2). This also further clarifies the trade-off between Oury and Tercieux (2012) and our paper. The trade-off is not that they allow indirect mechanism whereas we focus on direct revelation mechanisms. Our approach has more bite in the classical literature where messages are cheap talk. This enables us to study the robustness of the revelation principle (Assumption 3 reduces to Assumption 1 when applied to direct revelation mechanisms and truthful strategies being the equilibrium). The cost is that we need these kinds of “richness” assumptions to make any progress. Conversely, their approach needs no such richness assumption, but instead appeals to a vanishing cost of messages. This allows them to provide a full characterization of continuous implementation of a social choice function. In particular they show that continuous implementation is equivalent to rationalizable implementation of the social choice function in the baseline environment.

We should note that Theorem 5 only provides necessary but not sufficient conditions: the strict rationalizable correspondence need not be upper-hemicontinuous. Therefore we cannot conclude that a social choice function that satisfies this condition will be continuously implementable with respect to the product topology. Of course, we know from Oury and Tercieux (2012) that rationalizable implementability of the social choice function is sufficient. There is, therefore, a gap between the necessary and sufficient conditions in this setting. A full characterization appears out of reach.

## 5. An example: natural resource auctions

It may be useful at this stage to develop an example to help readers appreciate the implications of our results in a classical applied mechanism design setting.<sup>14</sup> To that end consider the following variant of a natural resource auction model.

There is a principal (e.g. the government) who wishes to auction a license to utilize a natural resource, e.g. a license to drill wells at a particular tract of land. The tract has an unknown quantity of oil  $q$ , which can take one of two values 0 and  $\bar{q}$  (i.e. the tract either contains no oil or a quantity  $\bar{q}$ ). Instead, agents see estimates. Each estimate  $e \in E = \{\underline{e}, \bar{e}\}$ .

The price of oil is normalized to 1. There are 2 competing buyers. The net value to buyer  $i$  of winning the license to operate the tract for a license fee of  $l$  is therefore  $q - l$ .

There is a finite set of feasible bids/payments  $B$ , so the set of alternatives the principal considers is  $A = \{1, 2\} \times B \times B$ , that is to say which buyer the license is allotted to and how much each buyer is charged.

*Baseline information structure* In baseline information structure, each quantity is equally likely. Estimates are assumed to be “correct” with probability  $\pi \in (\frac{1}{2}, 1)$ , i.e. a given estimate is  $\bar{e}$  in quantity state  $\bar{q}$ , and analogously for  $\underline{e}$  in the 0 quantity state. In the baseline information structure, it is common knowledge that each agent sees  $K$  estimates that are conditionally i.i.d. (conditional on quantity state). This baseline information structure defines a common-values setting.

In this baseline setting, the type of an agent can be summarized by a single number  $k \in \{0, 1, \dots, K\}$ , which is the number of high estimates they see in their  $K$  estimates.

*Social choice function* Define by  $q(k_i, k_j)$  the posterior expected quantity of oil in the tract when agent  $i$  sees  $k_i$  high estimates and agent  $j$  sees  $k_j$  high estimates. Suppose the space of feasible bids  $B$  is such that  $B = \{q(k, k) : k = 0, 1, \dots, K\}$ . Note that  $q(k, k)$  is strictly increasing in  $k$ . Consider the second-price auction where the strategy space of each agent is the set of bids  $B$ . Note that it is a strict Bayes-Nash equilibrium for each agent with type  $k$  to bid  $q(k, k)$  in this second-price auction (we omit the formal arguments for brevity—refer to e.g. Chapter 6 of Krishna (2009) for a textbook treatment).

The principal’s desired social choice function is the outcome of the second-price auction when agents play the BNE above, i.e. if agent 1 is of type  $k_1$  and agent 2 is of type  $k_2$ , then the agent with the larger type gets the good and pays the principal  $q(k, k)$  corresponding to the smaller type.

Now let us investigate two possible perturbations of this baseline information structure:

<sup>14</sup> We thank Muhamet Yildiz for suggesting this application.



- (1) In the first, it is common knowledge that each buyer has at most  $K$  estimates, but each buyer places some small probability that the other has fewer estimates. As motivation, consider that agents are allowed by law to only sample the tract at a maximum of  $K$  locations but consider the possibility that their competitor sampled fewer locations.
- (2) In the second, there is no upper bound on the number of estimates  $K$ . For example, each buyer considers that the competitor may have disobeyed the law and in fact sampled the tract in additional locations.

Intuitively, the former corresponds to perturbations that remain close to the baseline in the uniform-weak topology. The latter on the other hand is an e-mail game type structure that corresponds to perturbations in the product topology. For brevity we do not describe these type spaces explicitly. In such a setting, our results now have the following implications:

- (1) Theorem 1 tells us that the principal's desired social choice function cannot be truthfully continuously implemented with respect to perturbations of type (2). Formally, this is because the direct revelation mechanism corresponding that implements the social choice function has other rationalizable strategies for the agent—for example, it is an equilibrium for one agent to bid  $q(0, 0)$  and the other to bid  $q(K, K)$  independent of their types.
- (2) Theorem 2 tells us that the principal's desired social choice function can be truthfully continuously implemented with respect to perturbation (1)—because we have already argued that truth-telling is a strict BNE in the DRM.
- (3) Theorem 3 tells us that expanding the class of mechanisms does not expand the set of social choice functions that can be continuously implemented with respect to perturbations of type (1).

## 6. Related literature

There is a large, influential literature on the connection between higher-order beliefs and strategic behavior, beginning with the email game paper of Rubinstein (1989) and the subsequent global games paper of Carlsson and Van Damme (1993), too large to comprehensively cite here. Indeed, within this field there are now at least two influential approaches: the ex-ante approach of e.g. Kajii and Morris (1997), and the interim approach of Weinstein and Yildiz (2004) and Weinstein and Yildiz (2007). As we stated earlier, our approach borrows ideas from the latter.

There is also a large literature considering robustness in mechanism design. It bifurcates into “global” and “local” approaches.<sup>15</sup> In global approaches (see e.g. the pioneering works of Bergemann and Morris (2005); Chung and Ely (2007)) the planner has no information on the information structure (model) that will prevail among agents. The planner wishes to implement the social choice function on all models she considers possible. By contrast, in the local approach (see e.g. Chung and Ely (2003), Oury and Tercieux (2012), Jehiel et al. (2012) or Aghion et al. (2012)) the planner has some specific model in mind but is not entirely confident about it. The requirement therefore is analogously local, i.e. that the social choice function be implemented at types close to the initial model. This paper falls in the latter camp so we focus our discussion on related works in this vein.

<sup>15</sup> While we will not dwell on these, intermediate notions of robustness, where the principal rules out some possible beliefs among the agents, have also been recently formulated and characterized—see e.g. Ollár and Penta (2017).

The formulation of a “local” approach to robustness that we use in this paper was pioneered by Oury and Tercieux (2012). Our results have some counterparts to theirs. We therefore first discuss the connection to their paper before mentioning other work.

The biggest difference in setups is that we mainly consider implementation by “direct revelation mechanisms.” This assumption allows us tighter characterizations of (truthful) continuous implementation under the product topology. In the “forward” direction they consider the stronger desideratum of strict continuous implementation, and show that strict monotonicity of the social choice function is necessary for strict continuous implementation. To get a full characterization, and to study continuous implementation directly (as opposed to strict continuous implementation), they enrich the model to consider that sending various messages may involve small costs to the agents. By contrast, our assumptions allow us a full characterization without either (i.e. the strengthening of desideratum to strict continuous implementation, nor the possibility of costly messages). Another critical difference between our result and theirs is that our Theorem 1 is a characterization for the implementing DRM whereas their counterpart (Theorem 4) is a characterization of implementability (i.e., the mechanism that achieves rationalizable implementation is different from the mechanism that achieves continuous implementation in general (and also in their proof)).

They do not consider the uniform-weak topology but do hint at similar results in one direction (see, e.g., Footnote 16 of their paper). Our results on the uniform-weak topology thus both strengthen their results, and also constitute a key intermediate step to our characterization in the product topology.

Another closely related paper is that of Oury (2015), who characterizes continuous implementation as equivalent to full implementation in rationalizable strategies by introducing local payoff uncertainty of the planner. Assumption 1 in that paper embeds the set of states  $\Theta$  into a larger set of states  $\Theta^*$ , where these additional states allow to resolve indifferences.<sup>16</sup>

At a high level then, the difference between our approach and these two papers is that they consider general mechanisms, and obtain their characterization by extending the model (costly messages in the case of Oury and Tercieux (2012), additional states in the case of Oury (2015)). We instead cover only direct revelation mechanisms, and look at the robustness of a specific equilibrium (truthful equilibrium). Conversely our richness conditions (e.g. Assumption 1 or 2) can be verified directly within the benchmark model  $\bar{T}$  and our robustness exercise requires no extra payoff-relevant perturbation beyond what’s specified in the benchmark set of states  $\Theta$ .

A recent closely related paper that takes a different approach is Takahashi and Tercieux (2011): they study robust equilibrium *outcomes* rather than robust equilibrium *behaviors* (recall our discussion after Definition 1). Formally, they look at sequential games where there is almost common certainty of payoffs (for our purposes, “almost” refers to being close in the uniform-weak topology). The latter means that their results do not directly apply to our setting: Among other differences, our Theorem 3 requires the domain of the SCF to have a product structure, while almost common certainty implies the baseline type space diagonal. It is possible to construct an example of an SCF which is not truthfully continuously implementable with respect to  $d^{uw}$  but is implementable in a generic perfect-information extensive-form mechanism. It follows from Corollary 2 of their paper that the SCF is implementable in robust equilibrium outcomes

<sup>16</sup> In our notation, the definition of local payoff uncertainty is as follows (Assumption 1)—there is a baseline model  $\bar{T}$ , and the set of states of the world considered by types in the baseline is  $\Theta$ . However, the principal envisages a larger set of states  $\Theta^*$ , where  $\Theta \subseteq \Theta^*$  and for every agent  $i$ , alternative  $a$  and state  $\theta$  there exists a state  $\theta^*(\theta, a, i)$  such that  $u_i(a, \theta^*(\theta, a, i)) > u_i(a, \theta)$  and  $u_i(a', \theta^*(\theta, a, i)) = u_i(a', \theta)$  for any other  $a' \neq a$ .

in the sense of Takahashi and Tercieux (2011). However, we do not know whether the SCF is continuously implementable with respect to  $d^{uw}$ .

As we alluded to earlier, other papers have raised similar questions about “local” robust implementation. Chung and Ely (2003) ask about the possibility of (full) implementation in undominated Nash equilibrium while additionally requiring that Bayes-Nash equilibria of settings with arbitrarily small uncertainty also be close to the social choice function. They show that monotonicity of the social choice function is a necessary condition in their setting (while full implementation in undominated Nash equilibrium is possible for any social choice function under complete information). Aghion et al. (2012) consider subgame-perfect implementation under similar perturbations. Jehiel et al. (2012) get a negative result similar in interpretation to ours, but in a different setting, where the multi-dimensionality of agents’ signals drives the result. Postlewaite and Wettstein (1989) pursue the idea of a feasible, continuous function that achieves Walrasian outcomes in an exchange economy. Continuity is with respect to small perturbations of the initial endowments, as a substitute to modeling incentive constraints.

Our work is also connected to the literature on informational size beginning with McLean and Postlewaite (2002). These papers consider settings close to complete information, and argue what can be thought of as continuity results—when the state is approximate common knowledge, small transfers are sufficient to elicit the private information of agents. Most papers in this line consider settings with transfers, except Gerardi et al. (2009). Our results in the uniform-weak topology can be thought of as complementing their findings—both suggest that in settings with approximate common knowledge of the information structure, a desired social choice function may be implemented. While they consider richer settings, they also assume a common prior among agents that is known to the principal.

## Appendix A. Omitted proofs

**Theorem 1.** *Suppose that Assumption 1 holds. An SCF  $f$  is truthfully continuously implementable w.r.t.  $d^p$  by a DRM  $g$  if and only if it is implementable in unique rationalizable action profile in  $\tilde{g}$ .*

**Proof.** ( $\Leftarrow$ ): Let  $\mathcal{T}$  be a model with  $\mathcal{T} \supset \bar{\mathcal{T}}$ . Since  $T$  is countable and  $\bar{T}$  is finite, a standard fixed-point argument implies that there is a BNE  $\sigma$  in the game  $U(g, \mathcal{T})$ . Let  $\tilde{\sigma}$  be the strategy profile in  $\tilde{g}$  induced from  $\sigma$ , i.e., for each  $t \in T$ , we set  $\tilde{\sigma}(t)[\tilde{t}] = \sigma(t)[\tilde{t}]$  where  $\tilde{t}$  is the set of messages strategically equivalent to  $t$  in the DRM  $g$ . Since  $\sigma$  is a BNE in  $g$ , it follows that  $\tilde{\sigma}$  is also a BNE in  $\tilde{g}$ .

Since  $R^\infty(t, \tilde{g}) = \{\tilde{t}\}$  for  $t \in \bar{T}$ , by the upper hemicontinuity of the rationalizable correspondence  $R^\infty(\cdot, \tilde{g})$  (see, e.g., Theorem 2 of Dekel et al. (2006)), there is some  $\varepsilon > 0$  such that

$$d_i^p(t'_i, t_i) < \varepsilon \Rightarrow R_i^\infty(t'_i, \tilde{g}) = \{\tilde{t}_i\}$$

Since  $\tilde{\sigma}$  is a BNE in  $\tilde{g}$ , it follows that  $\tilde{\sigma}_i(t'_i) = \delta_{\tilde{t}_i}$  for any  $t'_i \in T_i$  with  $d_i^p(t'_i, t_i) < \varepsilon$ . Hence, for any  $t'_i \in \text{supp}\sigma_i(t_i)$ , we have that  $t'_i$  is equivalent to  $\tilde{t}_i$ . Define a strategy profile  $\sigma'$  in  $U(g, \mathcal{T})$  as

$$\sigma'_i(t'_i) \equiv \begin{cases} \delta_{\tilde{t}_i}, & \text{if } d_i^p(t'_i, t_i) < \varepsilon; \\ \sigma_i(t'_i), & \text{otherwise.} \end{cases}$$

Since  $\sigma$  is a BNE in  $U(g, \mathcal{T})$ , that  $\sigma'$  is also a BNE. Moreover,  $g(\tilde{t}) = f(t)$  for every  $t \in \bar{T}$  and by construction  $\sigma'$  also satisfies requirement (b) in Definition 1.

( $\Rightarrow$ ): Fix a DRM  $g$  that truthfully continuously implements  $f$  w.r.t.  $d^P$ . Since  $f$  is truthfully continuously implementable by  $g$  w.r.t.  $d^P$ ,  $f$  is truthfully continuously implementable by  $g$  w.r.t.  $d^{uw}$ . By Theorem 2 and Corollary 1,  $f$  is implementable in strict BNE in  $\tilde{g}$ .

The following lemma will be useful.

**Lemma 3.** *For each  $k \geq 1$  and  $\varepsilon \in (0, 1)$ , there is a countable model  $\mathcal{T}_{k,\varepsilon} \supset \overline{\mathcal{T}}$  such that  $T_{i,0,\varepsilon} \equiv \overline{T}_i$  and  $T_{i,k,\varepsilon} \equiv \left( \bigsqcup_{t_i \in \overline{T}_i} R_i^k(t_i, \tilde{g}) \right) \bigsqcup T_{i,k-1,\varepsilon}$ .*

Fix any BNE  $\tilde{\sigma}$  of the game  $U(\tilde{g}, \mathcal{T}_{k,\varepsilon})$  with  $\tilde{\sigma}(t) = \delta_{\tilde{t}}$  for every  $t \in \overline{\mathcal{T}}$ . This model has the property that for each type  $t_{i,k,\varepsilon}(\tilde{t}'_i, t_i)$  (the type in  $T_{i,k,\varepsilon}$  that corresponds to  $(\tilde{t}'_i, t_i)$  such that  $\tilde{t}'_i \in R_i^k(t_i, \tilde{g})$ ),

- (1)  $d_i^k(t_{i,k,\varepsilon}(\tilde{t}'_i, t_i), t_i^k) < \varepsilon$ ;
- (2)  $\tilde{\sigma}_i(t_{i,k,\varepsilon}(\tilde{t}'_i, t_i)) = \delta_{\tilde{t}'_i}$ .

This lemma appears a little convoluted but is at the heart of our proof. It constructs a countable model  $\mathcal{T}_{k,\varepsilon}$  with following property: Consider any Bayes Nash equilibrium  $\tilde{\sigma}$  of the game of incomplete information  $U(\tilde{g}, \mathcal{T}_{k,\varepsilon})$  with the property that types in  $\overline{\mathcal{T}}$  all report their type “truthfully.” In other words, each type  $t_i$  sends the reduced normal form message  $\tilde{t}_i$  in  $\tilde{g}$  corresponding to the equivalence class which the type  $t_i$  falls in. Further, consider any message  $\tilde{t}'_i \in R_i^k(t_i, \tilde{g})$ , i.e. any message that survives up to  $k$  rounds of iterated deletion of never best response in  $\tilde{g}$  for type  $t_i$  of player  $i$ .

The model  $\mathcal{T}_{k,\varepsilon}$  is constructed such that there exists a type of player  $i$ ,  $t_{i,k,\varepsilon}(\tilde{t}'_i, t_i)$  that is  $\varepsilon$ -close to  $t_i$  in their  $k$ -th-order beliefs; moreover, player  $i$  of type  $t_{i,k,\varepsilon}(\tilde{t}'_i, t_i)$  must play  $\tilde{t}'_i$  under the BNE  $\tilde{\sigma}$ .

Before we present the proof of Lemma 3, let us conclude the now routine proof of Theorem 1. Consider the countable model  $\mathcal{T}$  where  $T_i = \bigsqcup_{k=1}^{\infty} T_{i,k,\frac{1}{k}}$  and  $\mathcal{T}_{k,\frac{1}{k}}$  is given as in Lemma 3.

Since  $f$  is truthfully continuously implementable w.r.t.  $d^P$ , there is a BNE  $\sigma$  in the game  $U(g, \mathcal{T})$  such that requirements (a) and (b) in Definition 1 hold. Again,  $\sigma$  induces a BNE  $\tilde{\sigma}$  in  $\tilde{g}$ . Since  $\sigma(t) = \delta_{\tilde{t}}$  by requirement (b) of Definition 1, we have  $\tilde{\sigma}(t) = \delta_{\tilde{t}}$ .

Thus, it follows from Lemma 3 that for each  $\tilde{t}'_i \in R_i^{\infty}(t_i, \tilde{g})$ , for each  $k$ , there is a type  $t_{i,k,\frac{1}{k}}(\tilde{t}'_i, t_i) \in T_i$  such that

$$d_i^k(t_{i,k,\frac{1}{k}}(\tilde{t}'_i, t_i), t_i^k) \leq \frac{1}{k}, \quad (1)$$

and

$$\tilde{\sigma}_i(t_{i,k,\frac{1}{k}}(\tilde{t}'_i, t_i)) = \delta_{\tilde{t}'_i}.$$

It follows from (1) that  $d_i^P(t_{i,k,\frac{1}{k}}(\tilde{t}'_i, t_i), t_i) \rightarrow 0$ . Since  $\sigma$  satisfies requirement (b) in Definition 1, we know that it must be the case that  $\sigma_i(t_{i,k,\frac{1}{k}}(\tilde{t}'_i, t_i)) \rightarrow \delta_{\tilde{t}'_i}$ . Hence,  $\tilde{t}'_i = \tilde{t}_i$ .

Finally, since  $\tilde{t}'_i \in R_i^{\infty}(t_i, \tilde{g})$  is arbitrary, we conclude that  $\tilde{t}_i$  is the unique rationalizable message profile at  $t$  in  $\tilde{g}$ .  $\square$

**Proof of Lemma 3.** Formally, fix  $\varepsilon \in (0, 1)$  and we prove the claim by induction. First, the claim trivially holds for  $k = 0$ . Now we prove the claim for  $k \geq 1$ , assuming that it holds for  $k - 1$ .

Denote by  $\tilde{t}_i$  the messages of agent  $i$  in the reduced-form  $\tilde{g}$ . By definition (recall Definition 5), each  $\tilde{t}'_i \in R_i^k(t_i, \tilde{g})$  is a best response to some belief  $\mu_{-i} \in \Delta(\Theta \times \bar{T}_{-i} \times \tilde{T}_{-i})$  such that:

$$\text{marg}_{\Theta \times \bar{T}_{-i}} \mu_{-i} = \kappa_{t_i},$$

$$\text{and } \mu_{-i} \left( \left\{ (\theta, t_{-i}, \tilde{t}'_{-i}) : \tilde{t}'_{-i} \in R_{-i}^{k-1}(t_{-i}, \tilde{g}) \right\} \right) = 1.$$

By the induction hypothesis, there is a mapping  $\eta_{-i,k-1,\varepsilon}$  from each  $t_{-i} \in \bar{T}_{-i}$  and  $\tilde{t}'_{-i} \in R_{-i}^{k-1}(t_{-i}, \tilde{g})$  to a type  $t_{-i,k-1,\varepsilon}(\tilde{t}'_{-i}, t_{-i})$  such that (1) and (2) in Lemma 3 hold.

Since  $\tilde{t}'_i$  is in the reduced form  $\tilde{g}$  of the DRM  $g$ ,  $\tilde{t}'_i$  is the equivalent class which includes some  $t'_i \in \bar{T}_i$ . Then, define  $\kappa_{t_{i,k,\varepsilon}(\tilde{t}'_i, t_i)} \in \Delta(\Theta \times T_{-i,k,\varepsilon})$

$$\kappa_{t_{i,k,\varepsilon}(\tilde{t}'_i, t_i)} = (1 - \varepsilon) \left( \mu_{-i} \circ \eta_{-i,k-1,\varepsilon}^{-1} \right) + \varepsilon \bar{\kappa}_{t'_i}.$$

That is, with probability  $(1 - \varepsilon)$ , type  $t_{i,k,\varepsilon}(\tilde{t}'_i, t_i)$  believes that the state and the opponents' types follow a distribution that is induced from  $\mu_{-i}$  (in which each  $t_{-i,k-1,\varepsilon}(\tilde{t}'_{-i}, t_{-i})$  plays  $\tilde{\sigma}_{-i}(t_{-i,k-1,\varepsilon}(\tilde{t}'_{-i}, t_{-i})) = \delta_{\tilde{t}'_{-i}}$  by the induction hypothesis); with probability  $\varepsilon$ , type  $t_{i,k,\varepsilon}(\tilde{t}'_i, t_i)$  has the same belief as type  $t'_i$ . Since  $\tilde{t}'_i$  is a best response against  $\mu_{-i}$  and the strict/unique best response against  $\bar{\kappa}_{t'_i}$  in  $\tilde{g}$  (by Corollary 1), it follows that  $\tilde{\sigma}_i(t_{i,k,\varepsilon}(\tilde{t}'_i, t_i)) = \delta_{\tilde{t}'_i}$ . Moreover, since

$$d_{-i}^{k-1}(t_{-i,k-1,\varepsilon}(\tilde{t}'_{-i}, t_{-i}), t_{-i}^{k-1}) < \varepsilon,$$

we have that  $d_i^k(t_{i,k,\varepsilon}(\tilde{t}'_i, t_i), t_i^k) < \varepsilon$ .  $\square$

**Lemma 1.** *If an SCF  $f$  is truthfully continuously implementable by a DRM  $g$  with respect to  $d^{uw}$ , then, for every agent  $i$  and any pair of agent  $i$ 's types  $t_i$  and  $t'_i$ , either  $g$  strictly rewards truth-telling at  $t_i$  over  $t'_i$ ; or  $t_i$  always weakly dominates  $t'_i$  in  $g$ .*

**Proof.** Suppose that  $f$  is continuously implementable w.r.t.  $d^{uw}$  by a DRM  $g$ . Consider a model  $\mathcal{T} = (T, \kappa)$  defined as follows. Let

$$T_j = \bar{T}_j \bigsqcup_{(\theta', t_j, t'_{-j}) \in \Theta \times \bar{T}_j \times \bar{T}_{-j}} \bigsqcup_{n=1}^{\infty} \left\{ t_{j,n}^{(t_j, \theta', t'_{-j})} \right\},$$

where we set  $\kappa_{t_j} = \bar{\kappa}_{t_j}$  for every  $t_j \in \bar{T}_j$ ; moreover, let

$$\kappa_{t_{j,n}^{(t_j, \theta', t'_{-j})}} = \left( 1 - \frac{1}{n} \right) \bar{\kappa}_{t_j} + \frac{1}{n} \delta_{(\theta', t'_{-j})}, \forall n \in \mathbb{N}.$$

In words for every agent  $j$ , every baseline type  $t_j$  of that agent, every state  $\theta'$ , every baseline type profile  $t'_{-j}$  of the others, and every natural number  $n$ , there is a type in model  $\mathcal{T}$  which we denote  $t_{j,n}^{(t_j, \theta', t'_{-j})}$ . This type has a belief that is a convex combination with weight  $(1 - \frac{1}{n})$  on the belief of the original baseline type  $t_j$ , and weight  $\frac{1}{n}$  on the degenerate distribution corresponding to the state being  $\theta'$  and the others having type profile  $t'_{-j}$ .

Observe that since  $\kappa_{t_j} = \bar{\kappa}_{t_j}$  for every  $t_j \in \bar{T}_j$  for each agent  $j$ , each  $t_{-j} \in \bar{T}_{-j}$  has exactly the same profile of hierarchies of beliefs as its corresponding type in  $T_{-j}$ . Hence,

$$d_j^{uw} \left( \begin{pmatrix} t_j, \theta', t'_{-j} \\ t_{j,n} \end{pmatrix}, t_j \right) \rightarrow 0.$$

Suppose instead that for some agent  $i$ , and some pair of types  $t_i$  and  $t'_i$  in  $\bar{T}_i$ ,  $g$  neither strictly rewards truth-telling at  $t_i$  over  $t'_i$  nor does  $t_i$  always weakly dominate  $t'_i$  in  $g$ , i.e.,

$$\sum_{(\theta, t_{-i}) \in \Theta \times \bar{T}_{-i}} [u_i(g(t'_i, t_{-i}), \theta) - u_i(g(t_i, t_{-i}), \theta)] \bar{\kappa}_{t_i}[(\theta, t_{-i})] \geq 0 \quad (2)$$

and for some  $t'_{-i} \in \bar{T}_{-i}$  and  $\theta'$ ,

$$u_i(g(t'_i, t'_{-i}), \theta') - u_i(g(t_i, t'_{-i}), \theta') > 0. \quad (3)$$

Hence, under  $\sigma_{-i}$ , by reporting  $t'_i$ , agent  $i$  with type  $t_{i,n}^{(t_i, \theta', t'_{-i})}$  gets interim expected payoff equal to

$$\left(1 - \frac{1}{n}\right) \sum_{(\theta, t_{-i}) \in \Theta \times \bar{T}_{-i}} u_i(g(t'_i, t_{-i}), \theta) \bar{\kappa}_{t_i}[(\theta, t_{-i})] + \frac{1}{n} u_i(g(t'_i, t'_{-i}), \theta').$$

Then, by (2) and (3), for agent  $i$  with this type, reporting  $t'_i$  is strictly better than reporting  $t_i$  for all  $n$  large enough. But then it cannot be the case that  $\sigma_i \left( t_{i,n}^{(t_i, \theta', t'_{-i})} \right) \rightarrow \delta_{t_i}$ . This contradicts the supposition that  $g$  truthfully continuously implements  $f$ .  $\square$

**Theorem 2.** An SCF  $f$  is truthfully continuously implementable by a DRM  $g$  with respect to  $d^{uw}$  if and only if  $g(t) = f(t)$  for all  $t \in \bar{T}_0$  and, for every agent  $i$  and any pair  $t_i$  and  $t'_i$ , either  $g$  strictly rewards truth-telling at  $t_i$  over  $t'_i$ ; or  $t_i$  is strategically equivalent to  $t'_i$  for agent  $i$ .

**Proof.** ( $\Rightarrow$ ) Observe that when a DRM  $g$  strictly rewards truth-telling at  $t_i$  over  $t'_i$ , then  $t'_i$  cannot always weakly dominate  $t_i$ . Thus, it follows from Lemma 1 that if  $f$  is truthfully continuously implementable by a DRM  $g$ , then  $t'_i$  always weakly dominates  $t_i$  if and only if  $t_i$  always weakly dominates  $t'_i$ , i.e., they are strategically equivalent in the sense of Definition 3.

( $\Leftarrow$ ) Let  $g$  be a DRM such that for every agent  $i$  and any pair  $t_i$  and  $t'_i$ , either  $g$  strictly rewards truth-telling at  $t_i$  over  $t'_i$ ; or  $t_i$  is strategically equivalent to  $t'_i$  for agent  $i$ . Further suppose that  $g(t) = f(t)$  for all  $t \in \bar{T}_0$ . Fix a model  $\mathcal{T} = (T, \kappa)$  and we show that  $g$  truthfully continuously implements  $f$  with respect to  $d^{uw}$  by constructing a BNE  $\sigma$  in  $U(g, \mathcal{T})$ .

Since  $\bar{T}$  is nonredundant, it follows from Mertens and Zamir (1985) that  $\bar{T}$  can be embedded into  $T^*$  such that  $\bar{\kappa}_{t_i} = \kappa_i^*(t_i)$  for each  $t_i \in \bar{T}_i$ .<sup>17</sup> Then, for each  $i$ , and  $t_i$  and  $t'_i$  in  $\bar{T}_i$  such that  $g$  strictly rewards truth-telling, we have

$$\sum_{(\theta, t_{-i}) \in \Theta \times \bar{T}_{-i}} [u_i(g(t_i, t_{-i}), \theta) - u_i(g(t'_i, t_{-i}), \theta)] \kappa_i^*(t_i)[(\theta, t_{-i})] > 0. \quad (4)$$

<sup>17</sup> More precisely, denote by  $h_j(t_j) = (t_j^1, t_j^2, \dots)$  the hierarchy of beliefs of  $t_j$ . If  $\bar{T}$  is nonredundant, then  $h_j$  is bijection between  $\bar{T}_j$  and a subset of  $T_j^*$  for every  $j$ . Moreover, Mertens and Zamir (1985) show that  $\bar{\kappa}_{t_i} = \kappa_i^*(h_i(t_i)) \circ (I_\Theta \times h_{-i})^{-1}$  where  $I_\Theta$  is the identity mapping on  $\Theta$  and  $h_{-i} \equiv (h_j)_{j \neq i}$ . We abuse the notation here to write  $\bar{\kappa}_{t_i} = \kappa_i^*(t_i)$ .

Again, since  $\bar{T}$  is nonredundant,  $t_{-i}^{k-1} \neq t_{-i}'$  for some  $k$  whenever  $t_{-i}' \neq t_{-i}$  and  $t_{-i}', t_{-i} \in \bar{T}_{-i}$ . Since  $\bar{T}$  is finite, for any  $k$  sufficiently large,  $t_{-i}^{k-1} \neq t_{-i}'$  whenever  $t_{-i}' \neq t_{-i}$  and  $t_{-i}', t_{-i} \in \bar{T}_{-i}$ . Let  $\{(\theta, t_{-i})\}^{k-1, \varepsilon}$  be the  $\varepsilon$ -ball around  $(\theta, t_{-i})$  with respect to the distance  $\max\{d^0, d^1, \dots, d^{k-1}\}$ . Since  $\bar{T}$  is finite, we may further decrease  $\varepsilon$  such that  $\{(\theta, t_{-i})\}^{k-1, \varepsilon}$  and  $\{(\theta, t_{-i}')\}^{k-1, \varepsilon}$  are disjoint whenever  $t_{-i}' \neq t_{-i}$ .

Let  $\mathcal{T}' = (T', \kappa^*)$  be the model induced by  $\mathcal{T} = (T, \kappa)$  such that  $T' \subset T^*$  is also countable and each  $t_i'' \in T_i$  is mapped to its hierarchy of belief which we also denote by  $t_i'' \in T_i'$ .

Now consider any  $t_i \in \bar{T}_i$  and any type  $t_i'' \in T_i'$  such that  $d_i^{\text{uw}}(t_i'', t_i) < \varepsilon$ . We show that

$$|\kappa_i^*(t_i'') [\{(\theta, t_{-i})\}^{\infty, \varepsilon}] - \kappa_i^*(t_i) [(\theta, t_{-i})]| \leq \varepsilon \quad (5)$$

where  $\{(\theta, t_{-i})\}^{\infty, \varepsilon}$  denotes the  $(d_i^{\text{uw}}, \varepsilon)$ -ball around  $(\theta, t_{-i})$ . Since  $d_i^{\text{uw}}(t_i'', t_i) < \varepsilon$ , we have  $d_i^k(t_i'', t_i^k) < \varepsilon$  for every  $k$ . Define

$$\{(\theta, t_{-i})\}^{*, k-1, \varepsilon} \equiv \left\{ (\theta'', t_{-i}'') \in \Theta \times T_{-i}' : (\theta'', (t_{-i}')^{k-1}) \in \{(\theta, t_{-i})\}^{k-1, \varepsilon} \right\}.$$

Since  $d_i^k(t_i'', t_i^k) < \varepsilon$ , we have

$$\begin{aligned} \kappa_i^*(t_i'') [\{(\theta, t_{-i})\}^{*, k-1, \varepsilon}] &= (t_i'')^k [\{(\theta, t_{-i})\}^{k-1, \varepsilon}] \\ &\geq t_i^k \left[ (\theta, t_{-i}^{k-1}) \right] - \varepsilon \\ &= \kappa_i^*(t_i) [(\theta, t_{-i})] - \varepsilon, \end{aligned} \quad (6)$$

where the equalities follow from the construction of the homeomorphism  $\kappa_i^*$  in Mertens and Zamir (1985). Likewise,

$$\begin{aligned} &\kappa_i^*(t_i'') \left[ \bigcup_{t_{-i}' \in \bar{T}_{-i} : t_{-i}' \neq t_{-i}} \{(\theta, t_{-i}')\}^{*, k-1, \varepsilon} \right] \\ &= (t_i'')^k \left[ \bigcup_{t_{-i}' \in \bar{T}_{-i} : t_{-i}' \neq t_{-i}} \{(\theta, t_{-i}')\}^{k-1, \varepsilon} \right] \\ &\geq \sum_{t_{-i}' \in \bar{T}_{-i} : t_{-i}' \neq t_{-i}} t_i^k \left[ (\theta, t_{-i}'^{k-1}) \right] - \varepsilon \\ &= \sum_{t_{-i}' \in \bar{T}_{-i} : t_{-i}' \neq t_{-i}} \kappa_i^*(t_i) [(\theta, t_{-i}')] - \varepsilon. \end{aligned} \quad (7)$$

Since  $\{(\theta, t_{-i})\}^{k-1, \varepsilon}$  and  $\{(\theta, t_{-i}')\}^{*, k-1, \varepsilon}$  are disjoint whenever  $t_{-i}' \neq t_{-i}$ , we have

$$\begin{aligned} \kappa_i^*(t_i'') [\{(\theta, t_{-i})\}^{*, k-1, \varepsilon}] &\leq 1 - \kappa_i^*(t_i'') \left[ \bigcup_{t_{-i}' \in \bar{T}_{-i} : t_{-i}' \neq t_{-i}} \{(\theta, t_{-i}')\}^{*, k-1, \varepsilon} \right] \\ &\leq 1 - \sum_{t_{-i}' \in \bar{T}_{-i} : t_{-i}' \neq t_{-i}} \kappa_i^*(t_i) [(\theta, t_{-i}')] + \varepsilon \\ &= \kappa_i^*(t_i'') [(\theta, t_{-i})] + \varepsilon \end{aligned} \quad (8)$$



where the second inequality follows from (7). Since  $\{(\theta, t_{-i})\}^{k,\varepsilon} \downarrow \{(\theta, t_{-i})\}^{\infty,\varepsilon}$  as  $k \rightarrow \infty$ , we have

$$\kappa_i^*(t_i'') [\{(\theta, t_{-i})\}^{\infty,\varepsilon}] = \lim_{k \rightarrow \infty} \kappa_i^*(t_i'') [\{(\theta, t_{-i})\}^{k,\varepsilon}]. \quad (9)$$

Now that (6) and (8) hold for all  $k$  sufficiently large, the right-hand side of (9) belongs to  $[\kappa_i^*(t_i) [(\theta, t_{-i})] - \varepsilon, \kappa_i^*(t_i) [(\theta, t_{-i})] + \varepsilon]$ . Hence, (5) holds, as desired.

Since (4) and (5) hold and  $\bar{T}$  is finite, pick sufficiently small  $\varepsilon > 0$  such that for any  $t_i'' \in T_i'$  with  $d_i^{uw}(t_i'', t_i) < \varepsilon$ ,

$$\sum_{(\theta, t_{-i}) \in \Theta \times T_{-i}} [u_i(g(t_i, t_{-i}), \theta) - u_i(g(t_i', t_{-i}), \theta)] \kappa_i^*(t_i'') [\{(\theta, t_{-i})\}^{\infty,\varepsilon}] > \varepsilon D \quad (10)$$

where

$$D \equiv \max_{i, t, t', \tilde{\theta}} |u_i(g(t), \tilde{\theta}) - u_i(g(t'), \tilde{\theta})|.$$

Consider the agent normal-form of the game  $U(g, \mathcal{T}')$  with the restriction that  $t_i''$  in the  $(d_i^{uw}, \varepsilon)$ -ball around  $t_i$  must report  $t_i$ . Denote this game with restriction by  $\bar{U}(g, \mathcal{T}')$ .

Since  $T'$  is countable and  $\bar{T}$  is finite, a standard fixed-point argument implies that  $\bar{U}(g, \mathcal{T}')$  has a BNE  $\sigma$ . By construction of  $\bar{U}(g, \mathcal{T}')$ , for any sequence  $d^{uw}(t_n, t) \rightarrow 0$ , we have  $\sigma(t_n) = t$  for  $n$  large enough.

Furthermore,  $\sigma$  is a BNE in the game  $U(g, \mathcal{T}')$ . To see this, note that for any agent  $i$  in the  $\varepsilon$ -ball around  $t_i$ , given that all other agents  $-i$  in the  $\varepsilon$ -ball around  $(\theta, t_{-i})$  are reporting  $t_{-i}$ , the unique best response (modulo strategic equivalence) is to play  $t_i$ . This follows from (10). Then, by Proposition 7 of Friedenberg and Meier (2017),  $\sigma \circ h$  is a BNE in the game  $U(g, \mathcal{T})$  where  $h = (h_i)_{i \in I}$  with  $h_i : T_i \rightarrow T_i'$  mapping each  $t_i''$  in  $T_i$  to its hierarchy of beliefs. Therefore,  $g$  truthfully continuously implements  $f$  with respect to  $d^{uw}$ .  $\square$

**Lemma 2.** *If  $\bar{T}_0 = \bar{T}$  and  $f$  is continuously implementable w.r.t.  $d^{uw}$  by mechanism  $\mathcal{M} = (M, g)$ , then there is a pure-strategy BNE  $\bar{\sigma}$  in  $U(\mathcal{M}, \bar{T})$  such that for each agent  $i$ , each type  $t_i$  in  $\bar{T}_i$  and message  $m_i' \in M_i$ , either  $f$  strictly rewards  $\bar{\sigma}_i(t_i)$  over  $m_i'$ ; or  $\bar{\sigma}_i(t_i)$  always weakly dominates  $m_i'$  in BNE  $\bar{\sigma}$ .*

**Proof.** Suppose that  $f$  is continuously implementable w.r.t.  $d^{uw}$  by mechanism  $\mathcal{M} = (M, g)$ . Consider a model  $\mathcal{T} = (T, \kappa)$  defined as follows. Let

$$T_j = \bar{T}_j \sqcup \bigsqcup_{(\theta', t_j, t_{-j}') \in \bar{T}_j \times \Theta \times \bar{T}_{-j}} \bigsqcup_{n=1}^{\infty} \left\{ t_{j,n}^{(t_j, \theta', t_{-j}')} \right\},$$

where we set  $\kappa_{t_j} = \bar{\kappa}_{t_j}$  for every  $t_j \in \bar{T}_j$ ; moreover, let

$$\kappa_{t_{j,n}^{(t_j, \theta', t_{-j}')}} = \left(1 - \frac{1}{n}\right) \bar{\kappa}_{t_j} + \frac{1}{n} \delta_{(\theta', t_{-j}')}, \forall n \in \mathbb{N}.$$

It is straightforward to verify that  $d_j^{uw} \left( t_{j,n}^{(t_j, \theta', t_{-j}')} , t_j \right) \rightarrow 0$ . Since  $f$  is continuously implementable w.r.t.  $d^{uw}$  by  $\mathcal{M}$ , there is an equilibrium  $\sigma$  which continuously implements  $f$  in

$U(\mathcal{M}, \mathcal{T})$ . Since  $\sigma$  continuously implements  $f$  in  $U(\mathcal{M}, \mathcal{T})$ , we have (a)  $\sigma|_{\bar{T}}$  is a pure-strategy BNE in  $U(\mathcal{M}, \bar{T})$ ; (b)  $g(\sigma(t_n)) \rightarrow f(t)$  for any sequence of type profiles  $\{t_n\} \subset T$  and  $t \in \bar{T}$  with  $d^{\text{uw}}(t_n, t) \rightarrow 0$ . Since  $\bar{T}_0 = \bar{T}$ , it follows that

$$g(\sigma(t)) = f(t), \forall t \in \bar{T}, \quad (11)$$

$$g\left(\sigma_i\left(t_{i,n}^{(t_i, \theta', t'_{-i})}\right), \sigma_{-i}(t_{-i})\right) \rightarrow f(t_i, t_{-i}), \forall t_{-i} \in \bar{T}_{-i}. \quad (12)$$

Suppose to the contrary that for some agent  $i$ , the SCF  $f$  neither strictly rewards  $t_i$  over  $m'_i$ ; nor does  $t_i$  always weakly dominate  $m'_i$  in  $\sigma$  (or more precisely  $\bar{\sigma} \equiv \sigma|_{\bar{T}}$  in  $U(\mathcal{M}, \bar{T})$ ), i.e.,

$$\sum_{(\theta, t_{-i}) \in \Theta \times \bar{T}_{-i}} [u_i(g(m'_i, \sigma_{-i}(t_{-i})), \theta) - u_i(g(\sigma_i(t_i), \sigma_{-i}(t_{-i})), \theta)] \bar{\kappa}_{t_i}[(\theta, t_{-i})] \geq 0 \quad (13)$$

and for some  $t'_{-i}$  and  $\theta'$ ,

$$u_i(g(m'_i, \sigma_{-i}(t'_{-i})), \theta') - u_i(g(\sigma_i(t_i), \sigma_{-i}(t'_{-i})), \theta') > 0. \quad (14)$$

First, it follows from (13) and (14) that

$$\begin{aligned} & \left(1 - \frac{1}{n}\right) \sum_{(\theta, t_{-i}) \in \Theta \times \bar{T}_{-i}} u_i(g(m'_i, \sigma_{-i}(t_{-i})), \theta) \bar{\kappa}_{t_i}[(\theta, t_{-i})] \\ & + \frac{1}{n} u_i(g(m'_i, \sigma_{-i}(t'_{-i})), \theta') \\ & > \left(1 - \frac{1}{n}\right) \sum_{(\theta, t_{-i}) \in \Theta \times \bar{T}_{-i}} u_i(g(\sigma_i(t_i), \sigma_{-i}(t_{-i})), \theta) \bar{\kappa}_{t_i}[(\theta, t_{-i})] \\ & + \frac{1}{n} u_i(g(\sigma_i(t_i), \sigma_{-i}(t'_{-i})), \theta') \end{aligned} \quad (15)$$

where the left-hand side of (15) is the interim expected payoff of agent  $i$  with type  $t_{i,n}^{(t_i, \theta', t'_{-i})}$  under  $\sigma_{-i}$ , by reporting  $m'_i$ .

Second, by (12) and since there are only finitely many outcomes in  $f$ , for each  $t_{-i} \in \bar{T}_{-i}$ , there is some  $M_i^{t_{-i}} \subset M_i$  such that for any sufficiently large  $n$ ,

$$\begin{aligned} \sigma_i\left(t_{i,n}^{(t_i, \theta', t'_{-i})}\right) \left[M_i^{t_{-i}}\right] & \geq 1 - \frac{1}{2|\bar{T}_{-i}|}; \\ g(m_i, \sigma_{-i}(t_{-i})) & = f(t_i, t_{-i}) = g(\sigma_i(t_i), \sigma_{-i}(t_{-i})), \forall m_i \in M_i^{t_{-i}}. \end{aligned} \quad (16)$$

Then, since  $\bar{T}_{-i}$  is finite, it follows that for sufficiently large  $n$ , we have

$$\sigma_i\left(t_{i,n}^{(t_i, \theta', t'_{-i})}\right) \left[\bigcap_{t_{-i} \in \bar{T}_{-i}} M_i^{t_{-i}}\right] \geq \sum_{t_{-i} \in \bar{T}_{-i}} \sigma_i\left(t_{i,n}^{(t_i, \theta', t'_{-i})}\right) \left[M_i^{t_{-i}}\right] - (|\bar{T}_{-i}| - 1) > 0 \quad (17)$$

Finally, by (16), under  $\sigma_{-i}$ , by reporting  $m_i \in \bigcap_{t_{-i} \in \bar{T}_{-i}} M_i^{t_{-i}}$ , agent  $i$  of type  $t_{i,n}^{(t_i, \theta', t'_{-i})}$  gets the interim expected payoff equal to the right-hand side of (15). Hence, it follows from (15)

that agent  $i$  of type  $t_{i,n}^{(t_i, \theta', t'_{-i})}$  can profitably deviate by re-assigning the probability placed on  $\bigcap_{t_{-i} \in \bar{T}_{-i}} M_i^{t_{-i}}$  to  $m'_i$  instead. This contradicts to  $\sigma$  being a BNE.  $\square$

**Theorem 4.** Suppose that Assumption 3 holds for mechanism  $\mathcal{M}$ . Then,  $f$  is continuously implementable in  $d^p$  if and only if it is strictly continuously implementable in  $d^p$ .

**Proof.** Suppose that  $\mathcal{M}$  continuously implements  $f$  w.r.t.  $d^p$ . To prove that  $\tilde{\mathcal{M}}$  strictly continuously implements  $f$  w.r.t.  $d^p$ , consider any model  $\mathcal{T}' = (T', \kappa')$ . Denote by  $\mathcal{T}'' = (T'', \kappa'')$  the disjoint union of  $\mathcal{T}' = (T', \kappa')$  and the model  $\mathcal{T} = (T, \kappa)$  constructed in Lemma 2. Then, we must have some BNE  $\sigma$  which continuously implements  $f$  in  $U(\mathcal{M}, \mathcal{T}'')$  in  $d^p$  (and there by in  $U(\mathcal{M}, \mathcal{T})$  in  $d^{uw}$ ). It follows from Lemma 2 that  $\sigma|_{\bar{\mathcal{T}}}$  satisfies the property that for each agent  $i$ , each type  $t_i$  in  $\bar{T}_i$  and message  $m'_i \in M_i$ , we have either  $\sigma|_{\bar{\mathcal{T}}}$  strictly rewards  $t_i$  over  $m'_i$ ; or  $t_i$  always weakly dominates  $m'_i$  in BNE  $\sigma|_{\bar{\mathcal{T}}}$ . Since  $\bar{\mathcal{T}}$  has full support, if  $t_i$  always weakly dominates  $m'_i$  in  $\sigma|_{\bar{\mathcal{T}}}$  and  $\sigma_i|_{\bar{\mathcal{T}}}(t_i)$  is not strategically equivalent to  $m'_i$ , the message  $m'_i$  must yield strictly lower payoff than  $\sigma_i|_{\bar{\mathcal{T}}}(t_i)$  for type  $t_i$ . Hence, it follows from Assumption 3 that  $\sigma|_{\bar{\mathcal{T}}}$  is a strict BNE in  $U(\tilde{\mathcal{M}}, \bar{\mathcal{T}})$ . It follows that  $\sigma$  continuously implements  $f$  in  $U(\mathcal{M}, \mathcal{T}')$ . Hence,  $\tilde{\mathcal{M}}$  strictly continuously implements  $f$  w.r.t.  $d^p$ .  $\square$

**Theorem 5.** Suppose that Assumption 3 holds. An SCF  $f$  is continuously implementable w.r.t.  $d^p$  by a finite mechanism only if  $f$  is implementable in strictly rationalizable action profiles by a finite mechanism.

**Proof.** Since  $\mathcal{M} = (M, g)$  continuously implements  $f$  w.r.t.  $d^p$ , by Theorem 4, we may assume without loss of generality that  $\mathcal{M}$  strictly continuously implements  $f$  w.r.t.  $d^p$ . We start by proving the following key lemma.

**Lemma 4.** For each pure-strategy strict BNE  $\bar{\sigma}$  in  $U(\mathcal{M}, \bar{\mathcal{T}})$  and  $k \geq 0$ , there is a model  $\mathcal{T}_k^{\bar{\sigma}} \supset \bar{\mathcal{T}}$  such that  $T_{i,0}^{\bar{\sigma}} \equiv \bar{T}_i$  and

$$T_{i,k}^{\bar{\sigma}} \equiv \left( \bigsqcup_{t_i \in \bar{T}_i} W_i^k(t_i, \mathcal{M}) \right) \bigsqcup T_{i,k-1}^{\bar{\sigma}}.$$

Fix any BNE  $\sigma$  of the game  $U(\mathcal{M}, \mathcal{T}_k^{\bar{\sigma}})$  such that  $\sigma|_{\bar{\mathcal{T}}} = \bar{\sigma}$ . This model has the property that for each type  $t_{i,k}(m_i, t_i)$  (the type in  $T_{i,k}$  that corresponds to  $m_i \in W_i^k(t_i, \mathcal{M})$ ),

- (1)  $t_{i,k}^k(m_i, t_i) = t_i^k$ ;
- (2)  $\sigma_i(t_{i,k}(m_i, t_i)) = \delta_{m_i}$ .

Consider any message  $m_i \in W_i^k(t_i, \mathcal{M})$ , i.e. any message that survives up to  $k$  rounds of iterated deletion of never best response in  $\mathcal{M}$  for type  $t_i$  of player  $i$ . The model  $\mathcal{T}_k^{\bar{\sigma}}$  is constructed such that there exists a type of player  $i$ ,  $t_{i,k}(m_i, t_i)$  that has the same  $k$ -th-order beliefs; moreover, player  $i$  of type  $t_{i,k}(m_i, t_i)$  must play  $m_i$  under the BNE  $\sigma$ .

Before we present the proof of Lemma 4, let us conclude the now routine proof of Theorem 5. Consider the countable model  $\mathcal{T}$  where

$$T_i = \bigsqcup_{\bar{\sigma} \text{ is a pure-strategy strict BNE in } U(\mathcal{M}, \bar{\mathcal{T}})} \left( \bigsqcup_{k=0}^{\infty} T_{i,k}^{\bar{\sigma}} \right)$$

and  $T_{i,k}^{\bar{\sigma}}$  is given as in Lemma 4. Since  $\mathcal{M}$  strictly continuously implements  $f$  w.r.t.  $d^p$ , there is some BNE  $\sigma$  which strictly continuously implements  $f$  in  $U(\mathcal{M}, \mathcal{T})$ . Hence,  $\sigma|_{\bar{\mathcal{T}}}$  is a pure-strategy strict BNE in  $U(\mathcal{M}, \bar{\mathcal{T}})$ . It follows from Lemma 4 that for each  $k$  and each  $m_i \in W_i^k(t_i, \mathcal{M})$ , there is a type  $t_{i,k}^k(m_i, t_i) \in T_i$  such that

$$t_{i,k}^k(m_i, t_i) = t_i^k \quad (18)$$

and

$$\sigma_i(t_{i,k}(m_i, t_i)) = \delta_{m_i}.$$

It follows from (18) that  $d_i^p(t_{i,k}(m_i, t_i), t_i) \rightarrow 0$ . Since  $\mathcal{M}$  strictly continuously implements  $f$ , we know that it must be the case that  $g(\sigma(t_k(m, t))) \rightarrow f(t)$ . Since  $\sigma(t_k(m, t)) = m$ , it follows that  $g(m) = f(t)$  for every  $m \in W^\infty(t, \mathcal{M})$ .  $\square$

**Proof of Lemma 4.** First, since  $\sigma|_{\bar{\mathcal{T}}} = \bar{\sigma}$ , the claim trivially holds for  $k = 0$ . Now we prove the claim for  $k \geq 1$ , assuming that it holds for  $k - 1$ . By definition, each  $m_i \in W_i^k(t_i, \mathcal{M})$  is a strict best response to some belief  $\mu_{-i} \in \Delta(\Theta \times \bar{T}_{-i} \times M_{-i})$  such that  $\text{marg}_{\Theta \times \bar{T}_{-i}} \mu_{-i} = \kappa_{t_i}$  and  $\mu_{-i} \left( \left\{ (\theta, t_{-i}, m_{-i}) : m_{-i} \in W_{-i}^{k-1}(t_{-i}, \mathcal{M}) \right\} \right) = 1$ . By the induction hypothesis, there is a mapping  $\eta_{-i,k-1}$  from each  $t_{-i} \in \bar{T}_{-i}$  and  $m_{-i} \in W_{-i}^{k-1}(t_{-i}, \mathcal{M})$  to a type  $t_{-i,k-1}(m_{-i}, t_{-i})$  such that (1) and (2) in Lemma 4 holds. Define  $\kappa_{t_i,k}(m_i, t_i) \in \Delta(\Theta \times T_{-i,k}^{\bar{\sigma}})$  as

$$\kappa_{t_i,k}(m_i, t_i) = \mu_{-i} \circ \eta_{-i,k-1}^{-1}.$$

That is, type  $t_{i,k}(m_i, t_i)$  believes that the state and the opponents' types follow a distribution that is induced from  $\mu_{-i}$  (in which each  $t_{-i,k-1}(m_{-i}, t_{-i})$  plays  $m_{-i}$  in BNE  $\sigma$  by the induction hypothesis). Since  $m_i$  is a best response against  $\mu_{-i}$ , it follows that  $\sigma_i(t_{i,k}(m_i, t_i)) = \delta_{m_i}$ . Moreover, since  $t_{-i,k-1}^{k-1}(m_{-i}, t_{-i}) = t_{-i}^{k-1}$ , we have that  $t_{i,k}^k(m_i, t_i) = t_i^k$ .  $\square$

## References

- Aghion, P., Fudenberg, D., Holden, R., Kunitomo, T., Tercieux, O., 2012. Subgame perfect implementation under information perturbations. *Q. J. Econ.* 127 (4), 1843–1881.
- Bergemann, D., Morris, S., 2005. Robust mechanism design. *Econometrica* 73 (6), 1771–1813.
- Carlsson, H., Van Damme, E., 1993. Global games and equilibrium selection. *Econometrica* 61 (5), 989–1018.
- Chen, Y.-C., Di Tillio, A., Faingold, E., Xiong, S., 2010. Uniform topologies on types. *Theor. Econ.* 5 (3), 445–478.
- Chung, K.-S., Ely, J.C., 2003. Implementation with near-complete information. *Econometrica* 71 (3), 857–871.
- Chung, K.-S., Ely, J.C., 2007. Foundations of dominant-strategy mechanisms. *Rev. Econ. Stud.* 74 (2), 447–476.
- Dekel, E., Fudenberg, D., Morris, S., 2006. Topologies on types. *Theor. Econ.* 1, 275–309.
- Dekel, E., Fudenberg, D., Morris, S., 2007. Interim correlated rationalizability. *Theor. Econ.* 2 (1), 15–40.
- Friedenberg, A., Meier, M., 2017. The context of the game. *Econ. Theory* 63 (2), 347–386.
- Gerardi, D., McLean, R., Postlewaite, A., 2009. Aggregation of expert opinions. *Games Econ. Behav.* 65 (2), 339–371.
- Jehiel, P., Meyer-ter Vehn, M., Moldovanu, B., 2012. Locally robust implementation and its limits. *J. Econ. Theory* 147 (6), 2439–2452.
- Kajii, A., Morris, S., 1997. The robustness of equilibria to incomplete information. *Econometrica*, 1283–1309.
- Krishna, V., 2009. *Auction Theory*. Academic press.

- McLean, R., Postlewaite, A., 2002. Informational size and incentive compatibility. *Econometrica* 70 (6), 2421–2453.
- Mertens, J.-F., Zamir, S., 1985. Formulation of Bayesian analysis for games with incomplete information. *Int. J. Game Theory* 14 (1), 1–29.
- Monderer, D., Samet, D., 1989. Approximating common knowledge with common beliefs. *Games Econ. Behav.* 1, 170–190.
- Ollár, M., Penta, A., 2017. Full implementation and belief restrictions. *Am. Econ. Rev.* 107 (8), 2243–2277.
- Oury, M., 2015. Continuous implementation with local payoff uncertainty. *J. Econ. Theory* 159, 656–677.
- Oury, M., Tercieux, O., 2012. Continuous implementation. *Econometrica* 80 (4), 1605–1637.
- Postlewaite, A., Wettstein, D., 1989. Feasible and continuous implementation. *Rev. Econ. Stud.* 56 (4), 603–611.
- Rubinstein, A., 1989. The electronic mail game: strategic behavior under “almost common knowledge”. *Am. Econ. Rev.* 79 (3), 385–391.
- Satterthwaite, M.A., Sonnenschein, H., 1981. Strategy-proof allocation mechanisms at differentiable points. *Rev. Econ. Stud.* 48 (4), 587–597.
- Takahashi, S., Tercieux, O., 2011. Robust Equilibria in Sequential Games Under Almost Common Certainty of Payoffs. Discussion paper, mimeo.
- Weinstein, J., Yildiz, M., 2004. Finite-order implications of any equilibrium.
- Weinstein, J., Yildiz, M., 2007. A structure theorem for rationalizability with application to robust predictions of refinements. *Econometrica* 75 (2), 365–400.
- Weinstein, J., Yildiz, M., 2011. Sensitivity of equilibrium behavior to higher-order beliefs in nice games. *Games Econ. Behav.* 72 (1), 288–300.