

3

The Optimal Mechanism

We can use the decomposition implied by Abreu et al. [1] to write the Bellman equation that characterizes the frontier of equilibrium values that can be attained in this environment.

Define $\underline{v}$ as the expected value for a player when the *other* one always chooses the allocation; that is, $\underline{v} = E_{\theta} [u(\theta_i, \theta_{-i})]$. Analogously, define $\bar{v}$ as the players' payoff when his preferred action is always taken, $\bar{v} = E_{\theta} [u(\theta_i, \theta_i)]$. We assume that there exists a $\kappa > 0$ so that $\underline{w}_i - \kappa > \underline{v}$, $i = 1, 2$. In words, the payoff a player collects in case he exercises his outside option is bounded away from the payoff he gets when his opponent takes all decisions.

Let $D = \{\{a(\theta), w(\theta)\}_{\theta}\}$ be the set of all deterministic mechanisms. In order to induce enough convexity in our problem, we allow for arbitrary convex combinations of elements in D . We do so by allowing the mechanism to condition play on the realization of a public random device x which has a uniform distribution on $[0, 1]$. The set we consider is:

$$R = \left\{ \{a(\theta, x), w(\theta, x)\}_{\theta, x} \mid \text{for every } x, \{a(\theta, x), w(\theta, x)\}_{\theta} \in D \right\}$$

Then, the program of interest is:

$$V(v) = \sup_{\{\{a(\theta, x), w(\theta, x)\}_{\theta, x}\} \in R} E_{\theta, x} [(1 - \delta) u(a(\theta, x), \theta_n) + \delta V(w(\theta, x))] \quad (\text{P1})$$

subject to

$$E_{\theta, x} [(1 - \delta) u(a(\theta, x), \theta_2) + \delta w(\theta, x)] = v \quad (\text{PK})$$

$$E_{\theta_2, x} [(1 - \delta) u(a(\theta, x), \theta_1) + \delta V(w(\theta, x))] \geq E_{\theta_2, x} \left[(1 - \delta) u\left(a\left(\hat{\theta}_1, \theta_2, x\right), \theta_1\right) + \delta V\left(w\left(\hat{\theta}_1, \theta_2, x\right)\right) \right] \forall \theta_1, \hat{\theta}_1 \in \Theta \quad (\text{IC}_1)$$

$$E_{\theta_1, x} [(1 - \delta) u(a(\theta, x), \theta_2) + \delta w(\theta, x)] \geq E_{\theta_1, x} [(1 - \delta) u(a(\theta_1, \hat{\theta}_2, x), \theta_2) + \delta w(\theta_1, \hat{\theta}_2, x)] \forall \theta_2, \hat{\theta}_2 \in \Theta \quad (\text{IC}_2)$$

$$V(w(\theta, x)) \geq \underline{w}_1 \forall \theta \in \Theta^2 \text{ and } \forall x \in [0, 1] \quad (\text{IR}_1)$$

$$w(\theta, x) \geq \underline{w}_2 \forall \theta \in \Theta^2 \text{ and } \forall x \in [0, 1] \quad (\text{IR}_2)$$

The constraints are standard. The promise keeping (PK) constraint requires that, if agent two is promised discounted expected utility of v , the mechanism must choose an action $a(\cdot, \cdot)$ and continuation values $w(\cdot, \cdot)$ that deliver such promise. (IC₁) and (IC₂) state that, given a truthful report of the other agent, it must be optimal for agent i to report truthfully his preference shock.¹ Finally, the last two constraints are the participation constraints for agents 1 and 2, respectively.

Defining $\bar{w}_2 = V^{-1}(\underline{w}_1)$,² we can write the Participation Constraints as

$$w(\theta, x) \in [\underline{w}_2, \bar{w}_2] \forall \theta \in \Theta^2, x \in [0, 1] \quad (\text{IR}')$$

Early work (e.g. Casella [4], Jackson and Sonnenschein [8]) has shown that the repeated taking of joint actions allows for significant improvements over a one shot framework. The efficiency gains are attained by allowing the actions to be linked over time. A player who reports to have a more extreme preference shock – as measured by its distance from $\frac{1}{2}$ – is granted, relatively to a one shot case, more weight on the current action, relinquishing future decision power.

Define $a^*(\theta) = \arg\max_a E[u(a, \theta_1) + u(a, \theta_2)]$, to be the (ex-ante) Pareto efficient allocation, and let $v^{FB} = E[u(a^*(\theta, x), \theta_1) + u(a^*(\theta, x), \theta_2)]$ be the total surplus when action $a^*(\theta)$ is taken in all periods.

Under repeated decision taking, if either the Participation Constraint is ex-ante or players are forced to participate, v^{FB} can be arbitrarily approximated – but not attained – by equilibrium payoffs when players become patient. Carrasco and Fuchs [3], however, show that this can only be accomplished through the continuing variation in decision power. This variation, in turn, will necessarily lead to one of the players becoming the dictator: in the long run, one of the players will be promised $\bar{v}$.

In the current setting, this is not feasible because, whenever a player is promised sufficiently low continuation values, he will exercise his outside

¹ We make use of the Revelation Principle (Myerson [11]).

² Note that from the envelope theorem $V(\cdot)$ is strictly decreasing, thus it has an inverse.

option. As the mechanism cannot grant unbounded power to a player ex-post, it will not approximate efficiency ex-ante.

Theorem 1 (Inefficiency) *There exists $\epsilon > 0$ such that, for all $\delta \in [0, 1)$, the sum of the agents' payoffs is no larger than $v^{FB} - \epsilon$.*

Therefore, irrespective of how patient agents are, any feasible mechanism that satisfies the ex-post participation constraints will deliver outcomes that are bounded away from the efficient ones. Indeed, efficiency calls for intertemporal decisions to be linked: an agent who is given relatively more weight on a current decision has to relinquish future bargain power. The way through which the mechanism grants an agent a lower future bargain power is by promising him lower continuation values. The outside options place a lower bound on what a mechanism can promise to any single agent, impeding the mechanism to implement the efficient intertemporal trade of decision power.