

Time-consistent implementation in macroeconomic games

JEAN BARTHÉLEMY

Banque de France

ERIC MENGUS

Economics and Decision Sciences Department, HEC Paris and CEPR

The commitment ability of governments is neither infinite nor zero but intermediate. In this paper, we determine the commitment ability that a government needs to implement a unique equilibrium outcome and rule out self-fulfilling expectations. We show that, in a large class of static macroeconomic games, the government can obtain a unique equilibrium with any low level of commitment ability. We finally derive implications for models of bailouts and capital taxation.

KEYWORDS. Implementation, limited commitment, policy rules.

JEL CLASSIFICATION. C73, E58, E61, G28.

1. INTRODUCTION

In many macroeconomic games, private agents act anticipating future government's policy. Such an order of actions may lead to self-fulfilling expectations and multiple equilibria depending on the government's future policy response. As it is well known, in such situations, the government can commit to a rule that describes its future response to private actions to rule out undesired self-fulfilling expectations. This is, for example, one possible rationale for the “no bailout clause” in the EU Treaty (Article 125). The fear is that private bailout expectations fuel excessive risk taking, making bailouts ex post necessary. In contrast, the commitment not to bail out is supposed to steer agents to play the “good” equilibrium without socially costly bailouts and rules out “bad” equilibria with excessive risk taking.

However, governments are not fully committed to rules. Sometimes, depending on private actions, sticking to the rule is too costly for governments. Even if enshrined in a treaty, the “no bailout clause” was, for instance, not sufficient to completely rule out bailouts in the euro area (see, among others, [Gourinchas, Martin, and Messer \(2020\)](#)). In the extreme case of full discretion, governments may be tempted to respond to the

Jean Barthélemy: jean.barthelemy@banque-france.fr

Eric Mengus: mengus@hec.fr

This paper reflects the opinions of the authors and does not necessarily express the views of the Banque de France. We thank the editor and the referees, Kenza Benhima, Nicolas Coeurdacier, Hugues Dastarac, Alessandro Dovis, Stéphane Dupraz, Gaetano Gaballo, Stéphane Guibaud, Johan Hombert, Marie Laclau, Jean-Baptiste Michau, Franck Portier, Thomas Sargent, François Velde, Olivier Wang, and Pierre Yared as well as participants at many conferences and seminars for helpful comments. All remaining errors are ours.

private sector by confirming private sector's expectations, regardless of their past commitments. The resulting mutual feedback between private actions and the ex post policy response may then lead to multiple equilibria.¹

Still, governments have some ability to stick to rules. From simple speeches to more formal commitments such as contracts, laws, constitutions, treaties, or the delegations to independent authorities, commitments have in common to make future deviations costly, for example, from the simple embarrassment a policymaker may feel for breaking past promises² or the political costs of changing or breaching past legislations. Intuitively, the resulting limited commitment ability may allow governments to reach better outcomes. But, to what extent such limited commitment ability is sufficient to steer the private sector's expectations on a unique equilibrium? Alternatively, are multiple equilibria the unintended consequence of any limits to governments' commitment ability? If not, how rules should be designed to credibly prevent self-fulfilling private sector's expectations and ensure unique implementation?

To answer these questions, we consider a macroeconomic game between a large player—a government—and a large set of small agents—the private sector—that allows us to nest together the full spectrum of commitment abilities between full discretion and full commitment. We determine the conditions under which a government with a finite commitment ability can rule out equilibrium multiplicity. When it exists, we find the lowest commitment ability that the government needs to implement a unique equilibrium outcome—a situation that we refer to as (time-consistent) implementation. We describe the rules that can implement a unique equilibrium in a credible way, whenever possible. Finally, we derive implications for models of bailouts, inflation bias, and capital taxation.

More precisely, we add (Section 2) to a generic macroeconomic game an ex ante stage in which the government first commits to a reaction function that specifies policy responses as a function of every possible private sector's aggregate action, following Schelling (1960) and, more recently, Bassetto (2005). Then the private sector is competitive, that is, each private agent optimizes given the expectation of what other agents do and the government's future response. Finally, the government optimally selects its policy response. The government's payoff depends on the aggregate private action and on the policy response. The government also incurs a cost if the policy response deviates from the one that would follow the application of the reaction function. This cost measures the extent to which the government is bound by its commitments, and we will refer to this cost as the government's commitment ability.³ When this cost is zero, the government always chooses its ex post best response, regardless of its reaction function—this is the case of *full discretion*. In the other polar case, when this cost is infinite, the gov-

¹The literature review lists different literatures obtaining equilibrium multiplicity.

²As Woodford (2012) notes about the commitment to forward guidance announcements, "In practice, the most logical way to make such commitment achievable and credible is by publicly stating the commitment, in a way that is sufficiently unambiguous to make it embarrassing for policymakers to simply ignore the existence of the commitment when making decisions at a later time."

³We provide potential interpretations of this cost in Section 5.2.

ernment never deviates from its reaction function, which corresponds to the situation of *full commitment*.

Our main results are as follows. First, we derive the minimum commitment ability required for a government to implement its best time-consistent outcome for any macroeconomic game under rational expectations. Second, we show that, in many games with continuous action sets, an arbitrarily small commitment ability is enough for implementation. Limited commitment ability still has a strong impact on the design of credible rules, for example, a government can rule out bailout expectations by committing to a partial bailout close to, but below, the ex post optimal bailout. We show that this result, however, does not carry over with discrete action sets and we provide discussion that it may not be robust either to the introduction of imperfect information, fixed costs, or repeated interactions.

We first identify (Section 3) that, in static macroeconomic games, the critical parameter for implementation is what we dub the cost of controllability. Consider a private-sector action that the government wants to avoid. To rule it out, the government should pick a policy response that deters private agents from individually playing this action. There may exist many such responses, which may entail different costs for the government relative to its ex post best response. The cost of controllability is the maximum over all undesired private actions of the minimal cost of deterrence. The government implements a unique equilibrium outcome when its commitment ability exceeds this cost of controllability. Under this condition, the government can credibly commit to a reaction function from which it will not deviate ex post and that prevents any undesired outcome to form as an equilibrium. Finally, we show that a larger commitment ability always improves the best equilibrium outcome in static settings.

To our surprise, the cost of controllability boils down to zero in the static versions of the banks bailout from [Farhi and Tirole \(2012\)](#), the capital taxation problem (Section 4), and more generally to many games with continuous action sets. This implies that an arbitrarily small commitment ability is sufficient to implement a unique equilibrium. The main reason for this result is that, in these static models, private agents' marginal utilities are continuous functions of government responses, and the government can thus deter any undesired private actions by committing to a policy response arbitrarily close to its ex post best response. However, we show that, even with a zero cost of controllability, commitment ability still constrains the set of reaction functions that the government can use to obtain a unique equilibrium—and a larger commitment ability enlarges this set. With discrete action sets, however, the cost of controllability is generally strictly positive—deterring private actions requires bold actions by the government that are costly—and we provide an upper bound to this cost. We show that the closer the action sets are with a continuum, the more sensitive are private agents' payoffs with respect to policy, the less private agents' actions affect their marginal utility, the lower the cost of controllability.

An application of these results in the banks' bailout example is that the central bank can rule out bailout expectations with an arbitrarily low level of commitment ability. To this end, the central bank can commit to off-equilibrium partial bailouts that are only

slightly less generous than the bailout that is *ex post* optimal. These off-equilibrium partial bailouts are sufficient to make suboptimal any private actions anticipating a bailout. Thus, committing never to bail out is not necessary—on top of being time inconsistent for low commitment abilities. In the capital taxation example, inefficient equilibria under discretion, where capital is insufficiently accumulated and taxed at high rates, can be ruled out by off-equilibrium low-cost commitments, whereby the government commits to tax capital at a smaller rate compared to the *ex post* optimal rate.

Finally, we discuss the robustness of our results and the interpretations of the cost from deviating from the reaction (Section 5). In particular, we emphasize that our implementation result with continuous actions may not be robust to considering imperfect information, fixed costs or repeated interactions.⁴

Related literature First, our paper is connected to the literature on the time inconsistency of government policies, starting with Kydland and Prescott (1977) and Barro and Gordon (1983a). More recently, in a setting that is very close to ours, Dovis and Kirpalani (forthcoming) analyze the asymmetric information problem in which policymakers can either fully commit to rules or act under discretion. In contrast with their study, ours does not consider asymmetric information, and we focus on the ability of the government to implement a unique equilibrium outcome.

Within this literature, our paper is closer to the papers that, after Barro and Gordon (1983b), show that government's time-inconsistent policies lead to equilibrium multiplicity. Such a multiplicity of equilibria was obtained in multiple strands of the literature, in either static or dynamic settings.

In the literature on bailouts, the complementarity between private actions and bailout decisions is well known to produce multiple equilibria (see Schneider and Tornell (2004)). Farhi and Tirole (2012) build a static model in which the inability to commit not to bail out leads to additional inferior equilibria.⁵ This is a finding that is shared by Keister (2016) in a setting close to Ennis and Keister (2009) with the difference that some bailout is desirable even in the best equilibrium. In this literature, our paper is closely related to Philippon and Wang (2021), who also investigate “how much” commitment ability is needed to avoid bailout expectations, but in a setting in which policy can be made contingent on individual actions. Our finding that, at least in some models, the cost of controllability is small so that the coordination problem is easily solved, may be taken as a motive to focus on the effects of time inconsistency on the best equilibrium, as in Chari and Kehoe (2016). However, as we emphasize, this result is not robust to considering deviations from rationality or simply reputation forces, as in repeated settings, where, absent a large commitment ability, multiple equilibria may emerge.

⁴We provide a formal treatments of this statement in the working paper version of this paper Barthélemy and Mengus (2022).

⁵Notice that Farhi and Tirole (2012) emphasize that potential credibility losses may lead to a fixed cost for bailouts. However, they do not investigate the implications of such credibility losses for the required “amount” of commitment ability that would rule out equilibrium multiplicity.

In monetary economics, following Barro and Gordon (1983b), a literature has explored the conditions for the existence of multiple equilibria due to the central bank's time inconsistency—or “expectation traps” as dubbed by Chari, Christiano, and Eichenbaum (1998), who show how trigger strategies may lead to multiple equilibria, in a setting endogenizing both the public and the private sectors' behaviors. Time inconsistency also produces multiple equilibria even in the absence of trigger strategies (see Albanesi, Chari, and Christiano (2003), King and Wolman (2004), Armenter (2008)). In this paper, we abstract from trigger strategies and we refer the interested reader to the working paper version (Barthélemy and Mengus (2022)) for the analysis of implementation in repeated settings.

In monetary economics, the literature on monetary rules is also confronted to the presence of multiple equilibria. Here, the government or the central bank is able to commit to rules but, depending on its features, the rule itself may not be sufficient to prevent multiple equilibria to form (see Sargent and Wallace (1975), Taylor (1999), Clarida, Galí, and Gertler (2000), Loisel (2009), Atkeson, Chari, and Kehoe (2010), Hall and Reis (2016), among others). Our paper is in the tradition of this literature, which emphasized that rules should be state-contingent or, even, sophisticated—that is, history dependent. This literature has already highlighted constraints on off-equilibrium policy actions—they should be at least feasible as emphasized by Bassetto (2005)—or, they allow for a continuation of an equilibrium as in Atkeson, Chari, and Kehoe (2010), where the off-equilibrium central bank action is to keep the quantity of money constant forever. Historically, a key motive for introducing rules was to solve time-consistency issues as emphasized by Kydland and Prescott (1977). Such a motive has also led to investigate principal-agent approaches, delegations, and contract theory in monetary economics (see Rogoff (1985), Walsh (1995), Jensen (1997), Halac and Yared (forthcoming), among others). Our contribution with respect to this literature is to allow the government to deviate from its commitments, consistently with Bilbiie (2011) and Cochrane (2011). This has two implications. First, we consider the optimal design of commitments as part of the “game” played between the government and the private sector. This leads us to consider and model the government's incentives. Second, and more importantly, we investigate how limited commitment is not only a constraint on the design of rules but can also simply prevent the government to obtain a unique equilibrium.

The focus of the literature on taxation is usually more on the best equilibrium that can be sustained; see the recent contribution of Halac and Yared (2022). However, multiple equilibria may still emerge in frameworks such as those of Chari and Kehoe (1990) or Bassetto and Phelan (2008). In this literature, our paper is more closely connected to Farhi, Sleet, Werning, and Yeltekin (2012), who first investigate time-consistent capital taxation in a static model with an exogenous commitment ability and then consider the repeated setting to endogenize commitment ability. In contrast to their approach, we not only look at the best equilibrium outcome, but we investigate the full set of equilibria.

Finally, a literature on “loose commitment” following Debortoli and Nunes (2010) also introduces limited commitment ability, from discretion to full commitment, and

studies fiscal and monetary policy. Yet, their main focus is, so far, not on equilibrium multiplicity.

2. THE ENVIRONMENT

In this section, we describe the environment. The government first commits to a reaction function; then a continuum of atomistic private agents make decisions; finally, the government acts. The government incurs a welfare cost if its ex post action deviates from the reaction function it has committed to. The cost allows us to obtain a continuum of commitment abilities between full discretion and full commitment. We then define what we mean by a coordination problem and implementation in this environment.

Consider an economy populated by a continuum of identical private agents and a government. There are three stages. First, the government commits to a reaction function, $\bar{y}$, that maps any aggregate private-sector action $x \in X$ to a policy action $\bar{y}(x) \in Y$, where X and Y are compact intervals of $\mathbb{R}$. Second, each private agent chooses an action $\xi \in A_X \subseteq X$, where A_X is a closed subset of X . The average private action x is in the convex hull of A_X , contained in X . Finally, the government chooses an action $y \in A_Y \subseteq Y$, where A_Y is a closed subset of Y . As with feasibility conditions, private decisions constrain government actions, so any government action must belong to a nonempty closed subset $D(x) \subseteq A_Y$ that depends on the average private action x . Finally, we focus on pure strategies.

Competitive outcome For a given allocation (x, y) , the payoff of a private agent to play $\xi \in X$ is $u(\xi, x, y)$, where u is strictly concave and twice continuously differentiable in ξ . We define a competitive outcome as follows.

DEFINITION 1 (Competitive outcome). A competitive outcome is an allocation $(x, y) \in X \times A_Y$ such that

$$y \in D(x), \quad (1)$$

$$x \in \arg \max_{\xi \in A_X} u(\xi, x, y). \quad (2)$$

We denote by C the set of competitive outcomes.

Condition (1) requires that the government action y is feasible given average private action x , and Condition (2) requires that, given the allocation (x, y) , it is (weakly) optimal for any individual to set $\xi = x$. We focus on symmetric competitive outcome in which all private agents play the same action.⁶ Notice that by definition $x \in A_X$.

Finally, to avoid making the implementation problem trivial, we assume that the government cannot punish individual deviations directly. Thus, we assume that only the aggregate private outcome x , not individual decisions ξ , is public information.

⁶When the feasible set is the whole set, $A_X = X$, this restriction to symmetric outcome is without loss of generality as there is always a unique private best response ξ for any allocation (x, y) . When the feasible set A_X is not a segment, this restriction excludes heterogeneous private agents decision.

Government Before any action, the government commits to a reaction function $\bar{y}$. Such a reaction function corresponds to the commitment to take the action $\bar{y}(x)$ if the average private action is x . We assume that the government commits only to feasible actions; that is, for any $x \in X$, $\bar{y}(x) \in D(x)$.⁷ We denote by $Y(X)$ this set of functions.

The government cannot fully commit to the actions implied by its reaction function and can renege on its past commitments. However, if the government plays an action inconsistent with its reaction function, $y \neq \bar{y}(x)$, it incurs a cost $\kappa > 0$. κ measures the government value of sticking to a promise and will be referred to as the *commitment ability*. For a given reaction function $\bar{y}$ and an average private-sector action $x \in X$, the government selects an action $y \in D(x)$ by maximizing

$$r(\bar{y}, x, y, \kappa) = \begin{cases} w(x, y), & \text{if } y = \bar{y}(x), \\ w(x, y) - \kappa, & \text{otherwise,} \end{cases} \quad (3)$$

where w is a strictly concave and twice differentiable function in y .

The government selects its reaction function so as to maximize the ex ante payoff $\bar{r}(\bar{y}, x, y)$ defined as $r(\bar{y}, x, y)$, but where w is replaced by $\bar{w}$:

$$\bar{r}(\bar{y}, x, y, \kappa) = \begin{cases} \bar{w}(x, y), & \text{if } y = \bar{y}(x), \\ \bar{w}(x, y) - \kappa, & \text{otherwise,} \end{cases} \quad (4)$$

with $\bar{w}$, a strictly concave and twice differentiable function in y . By selecting the reaction function $\bar{y}$, the government affects its ex post incentives through the commitment ability κ . When $\kappa = 0$, the reaction function $\bar{y}$ is immaterial.

Notice that when $\bar{w}(x, y) = w(x, y) = u(x, x, y)$, aside from the reneging cost, the government is benevolent but disregards individual deviations. Considering a different payoff function for the ex ante choices by the government will be useful in some examples.

Timing and equilibrium An equilibrium is characterized by three strategies. They are, in chronological order: the reaction function $\bar{y} \in Y(X)$; the private-sector strategy $\sigma^h : Y(X) \rightarrow X$ that maps reaction function $\bar{y}$ into aggregate private-sector action $x = \sigma^h(\bar{y})$; and the government strategy σ^g specifies the government action $y = \sigma^g(\bar{y}, x) \in D(x)$ given the private-sector action $x \in X$ and the reaction function $\bar{y} \in Y(X)$. Figure 1 summarizes the timing of the game.

To define an equilibrium, we proceed in two stages. First, we define, if it exists, the continuation of an equilibrium given a reaction function. Second, we define the equilibrium itself. We do this as the continuation of an equilibrium may not necessarily form after any reaction function: the discontinuity in payoffs due to the cost κ may lead

⁷Under full commitment ability, committing to an unfeasible action would lead to a violation of a feasibility constraint. See [Bassetto \(2005\)](#) for further discussion about the role of feasibility constraints in implementation problem. In the absence of commitment, the government can still select a feasible action, despite the commitment to an unfeasible action. However, such a commitment would have no effects on ex post incentives to select one particular feasible action over another, and there is then no loss of generality in our assumption.

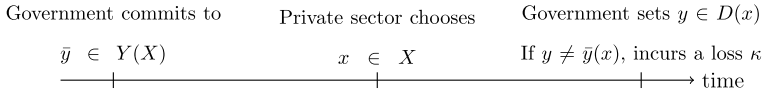


FIGURE 1. Timing of the game.

the government to favor ex post actions that are not compatible with any competitive outcome.

Let us first define the set of strategies for the private sector and the government's ex post action that allow for a continuation of an equilibrium after a reaction function

DEFINITION 2. For a given reaction function $\bar{\eta} \in Y(X)$, the strategies $\{h, g\}$ forms a **continuation of an equilibrium** when:

- (i) $(h(\bar{\eta}), g(\bar{\eta}, h(\bar{\eta})))$ is a competitive outcome; and
- (ii) for all $x \in X$, for all $\eta \in D(x)$, $r(\bar{\eta}, x, g(\bar{\eta}, x), \kappa) \geq r(\bar{\eta}, x, \eta, \kappa)$.

Conditions (i) and (ii) require that, given a reaction function $\bar{\eta}$, the strategies (h, g) constitute an equilibrium as in the standard models presented in [Stokey \(1991\)](#) or in Chapter 21 in [Ljungqvist and Sargent \(2018\)](#). The only noticeable difference is that the reaction function matters for condition (ii) through the cost κ in r , which depends on the reaction function $\bar{\eta}$.⁸

Let us denote by $\mathcal{CE}(\bar{\eta})$ the set of strategies $\{h, g\}$ that form the continuation of an equilibrium given a reaction function $\bar{\eta} \in Y(X)$. As we mentioned, $\mathcal{CE}(\bar{\eta})$ may well be empty for some reaction functions and we then denote by $\hat{Y}(X)$ the subset of $Y(X)$ so that $\mathcal{CE}(\bar{\eta})$ is not empty.

We then define an equilibrium as follows.

DEFINITION 3. A (subgame perfect) equilibrium is a reaction function $\bar{y}$, a private sector strategy σ^h , and a government strategy σ^g such that, for all $\bar{\eta} \in \hat{Y}(X)$:

- (i) $(\sigma^h, \sigma^g) \in \mathcal{CE}(\bar{\eta})$; and
- (ii) $\bar{r}(\bar{y}, \sigma^h(\bar{y}), \sigma^g(\bar{y}, x), \kappa) \geq \bar{r}(\bar{\eta}, \sigma^h(\bar{\eta}), \sigma^g(\bar{\eta}, \sigma^h(\bar{\eta})), \kappa)$.

An equilibrium is a reaction function and strategies for the private sector and the government so that, given the reaction function, the strategies form the continuation of an equilibrium (Condition (i)) and the reaction function leads to the best ex ante payoff for the government (Condition (ii)). Both conditions only require that, at each node of

⁸Notice that the government cannot play unfeasible actions, as we always require (both in- and off-equilibrium) that $y \in D(x)$ with x the action played by private agents. As a result, and as in [Bassetto \(2005\)](#), the government cannot rule out a competitive outcome by playing (or promising to play) unfeasible actions in our setting.

the game, actions are optimally selected as is the case in subgame perfect equilibria. Several additional comments are in order.

First, as we mentioned, the continuation of an equilibrium may not form after any reaction function. We do not consider this situation when assessing (ii) in the Definition 3, implicitly putting an arbitrarily low payoff to the absence of a continuation of an equilibrium.

Second, it is crucial in this definition that the private sector should play consistently with a competitive outcome after any reaction function, at least when possible $-(\sigma^h, \sigma^g) \in \mathcal{CE}(\bar{\eta})$. Individual optimality effectively puts a constraint on the private-sector strategy σ^h that has to react to the government's reaction function. Otherwise, without this constraint, the government would not be able to induce the private sector to change its behavior, as the private sector would be somehow able to “commit” to deviate from the competitive outcome to prevent the government from making certain commitments.

Finally, we are interested in the set of *equilibrium outcomes*—and not necessarily the precise strategy profile that leads to such an equilibrium outcome—and the properties of this set as a function of the commitment ability parametrized by κ . More formally, we have the following.

DEFINITION 4. Let us consider an equilibrium $(\bar{y}, \sigma^h, \sigma^g)$. Denoting by $x = \sigma^h(\bar{y})$ and $y = \sigma^g(\bar{y}, x)$, the resulting allocation $(x, y) \in X \times Y$ is an **equilibrium outcome**.

$\Theta(\kappa)$ denotes the set of such equilibrium outcomes and $v(\kappa)$ denotes the set of equilibrium (ex ante) payoffs for the government.

In particular, an allocation $(x, y) \in \Theta(0)$ is a **Nash outcome**.

One of our first tasks in the next subsection will be to move from Definition 4 to characterizations of the set of equilibrium outcomes. In what follows, we make the following assumption.

ASSUMPTION 1. *There exists at least one Nash outcome— $\Theta(0)$ is not empty.*

As we will see, this assumption is necessary for the set of equilibrium outcomes, $\Theta(\kappa)$, to be nonempty for any $\kappa > 0$.

Coordination problems and implementation We can define what we shall refer to as a coordination problem using the set of equilibrium outcomes $\Theta(0)$.

DEFINITION 5. The government faces a **coordination problem** when there exist multiple equilibrium outcomes under discretion; that is, $\Theta(0)$ contains more than one equilibrium outcome.

Given a commitment ability κ , the government **implements** an allocation (x, y) if the set of equilibrium outcomes satisfies $\Theta(\kappa) = \{(x, y)\}$.

A coordination problem is solved for a commitment ability κ when $\Theta(\kappa)$ is a singleton, in which case we also talk about implementation. The minimum for κ such that

$\Theta(\kappa)$ is a singleton is the measure of the commitment ability required to solve the coordination problem. Notice that a coordination problem may be immaterial for the government when the different equilibrium outcomes lead to the same ex ante payoff for the government. In contrast, a coordination problem is payoff-relevant for the government when $\inf v(0) < \sup v(0)$.

Ramsey allocation Finally, an allocation of interest is the Ramsey allocation, which we denote by (x^R, y^R) . It corresponds to one of best competitive outcomes given the ex ante objective function of the government—assuming that the government does not incur the reneging cost, κ :⁹

$$\bar{w}(x^R, y^R) = \max_{(x, y) \in C} \bar{w}(x, y). \quad (5)$$

This allocation coincides with the standard definition of Ramsey allocation when the government is benevolent ($\bar{w} = u$).

Notice that the Ramsey allocation (x^R, y^R) is not necessarily in the set of equilibrium outcomes $\Theta(\kappa)$, as playing y^R after x^R may be time inconsistent for the government. In contrast, when (x^R, y^R) belongs to the set of equilibrium outcomes—the government optimally plays y^R after the private sector plays x^R given the reaction function—the Ramsey allocation may not be the unique equilibrium outcome.¹⁰ Only in the case in which (x^R, y^R) is the only equilibrium outcome will we say that the government can implement the Ramsey allocation. This is obviously the ideal situation from the government's ex ante perspective.

3. IMPLEMENTATION UNDER LIMITED COMMITMENT ABILITY

In this section, we investigate the equilibrium set as a function of the commitment ability κ along the continuum $[0, \infty)$. Our first result is that implementation depends only on three simple objects: first, the best time-consistent competitive outcome; second, the *controllability*, which is the fact that the set of policy actions that deter private agents from expecting an inferior outcome is nonempty; third, the cost of playing such actions—the *cost of controllability*. When this latter cost is lower than the commitment ability, the government can solve the coordination problem and implement the best time-consistent competitive outcome. We then show that, under mild conditions, this cost of controllability is zero when the action set is continuous.

We start by defining the best time-consistent competitive outcome, controllability, and the cost of controllability.

Best time-consistent competitive outcome Let $(x^\kappa, y^\kappa) \in C$ be a competitive outcome that maximizes the government's ex ante welfare, under the constraint that the government prefers playing y^κ instead of deviating and paying the cost κ —that is playing y^κ is

⁹Note that multiple allocations may lead to the highest payoff for the government, because $\mathcal{A}_Y$ is not necessarily continuous.

¹⁰As Chari and Kehoe (2016) note, in this case, the Ramsey allocation is only *weakly* implemented.

time consistent after x^κ . Formally,

$$\max_{(x,y) \in C} \bar{w}(x, y) \quad (6)$$

$$\text{such that } w(x, y) \geq w(x, \eta) - \kappa, \forall \eta \in D(x). \quad (7)$$

As we will show, the best time-consistent competitive outcome is always the best equilibrium outcome.¹¹

Notice that (x^κ, y^κ) coincides with the Ramsey allocation (x^R, y^R) when $\kappa \geq \bar{\kappa}$, where $\bar{\kappa} \in \mathbb{R}^+$ is the lowest cost such that the constraint (7) is not binding for the Ramsey allocation:

$$\bar{\kappa} = \max_{\eta \in D(x^R)} w(x^R, \eta) - w(x^R, y^R). \quad (8)$$

The threshold $\bar{\kappa}$ measures the government's temptation to deviate from the Ramsey policy y^R when the private sector has played x^R .

Controllability What should the government commit to in order to control the private sector and force it to play the action consistent with the best constrained outcome x^κ ? In principle, private agents can play any action x that is consistent with a competitive outcome $(x, y) \in C$. To rule out a given private action x , the government should then commit to and stick to an action y such that (x, y) is not a competitive outcome. By definition, this means that the action y should make private agents better off playing $\xi \neq x$ if the aggregate private action is $x \neq x^\kappa$. Formally, for any $x \in A_X$, the government should select an action in the set:¹²

$$\mathcal{Y}(x) = \{y \in D(x) \mid \exists \xi \in A_X, u(\xi, x, y) > u(x, x, y)\}.$$

For any $x \in A_X \setminus \{x^\kappa\}$, the government *can* discourage private agents from playing x if this set $\mathcal{Y}(x)$ is nonempty.

ASSUMPTION 2 (Controllability). *For any $x \in A_X \setminus \{x^\kappa\}$, the set $\mathcal{Y}(x)$ is nonempty.*

Under this assumption, the government *can* control the private sector's actions and force it to play x^κ . But even if it *can*, it may be unable to credibly commit to all these policy actions because such actions are (ex post) costly. Take an aggregate private action $x \neq x^\kappa$. To deter that action, the government has to commit to an action $y \in \mathcal{Y}(x)$. Playing that action is costly to the extent that it is not ex post optimal: $y \notin \arg \max_{y \in D(x)} w(x, y)$.

¹¹ (x^κ, y^κ) exists. (x^0, y^0) exists, as $\Theta(0)$ is not empty following Assumption 1. As a result, the set of allocation $(x, y) \in C$ satisfying the constraint (7) is not empty. As this set is a compact set and $\bar{w}$ a continuous function, the optimization problem admits at least one solution. When multiple outcomes satisfy this problem, albeit this indeterminacy is not payoff relevant, the analysis carries over by arbitrarily selecting one of these outcomes.

¹²For any $x \notin A_X$, we already know that private agent will deviate from playing $\xi = x$ because x is not in the feasible set, and hence does not belong to any competitive outcome.

The resulting cost is then $w(x, y^*(x)) - w(x, y)$, with $y^*(x) \in \arg \max_{y \in D(x)} w(x, y)$.¹³ Naturally, this cost is 0 when playing one of the ex post optimal actions is sufficient to prevent private agents from playing x , that is, there exists $y^*(x) \in \arg \max_{y \in D(x)} w(x, y)$ such that $y^*(x) \in \mathcal{Y}(x)$. When the commitment ability κ is larger than this cost, the government will stick to its commitment and play y , and the expectation of such an action will deter agents from playing x .

Of course, for each $x \in A_X \setminus \{x^\kappa\}$, the government selects the action that minimizes its cost so that we can consider $\inf_{y \in \mathcal{Y}(x)} w(x, y^*(x)) - w(x, y)$. Finally, the government has to find such actions for all $x \in A_X \setminus \{x^\kappa\}$, so that the cost to deter any action that differs from x^κ is the following.

DEFINITION 6 (Cost of controllability). Let $\rho \geq 0$ be the cost of controllability with

$$\rho \equiv \sup_{x \in A_X \setminus \{x^\kappa\}} \inf_{y \in \mathcal{Y}(x)} [w(x, y^*(x)) - w(x, y)], \quad (9)$$

where $y^*(x) \in \arg \max_{y \in D(x)} w(x, y)$ is an ex post best action of the government with $\kappa = 0$.

The cost of controllability refers to the maximum cost that the government has to tolerate ex post to control the private sector. The cost of controllability is well-defined under Assumption 2 as, otherwise, no policy action may exist to deter the private sector from playing some private-sector action $x \neq x^\kappa$. The controllability assumption also implies that the cost of controllability is finite ($\rho < \infty$).¹⁴

The equilibrium set as a function of commitment ability We can now describe the equilibrium set $\Theta(\kappa)$ using the best time-consistent competitive outcome (x^κ, y^κ) and the cost of controllability ρ .

PROPOSITION 1. Under Assumption 2, the equilibrium set $\Theta(\kappa)$ is such that:

- (i) **Best time-consistent competitive outcome:** $(x^\kappa, y^\kappa) \in \Theta(\kappa)$. The set of equilibrium payoffs for the government $v(\kappa)$ is a compact set and $\bar{w}(x^\kappa, y^\kappa) = \max v(\kappa)$.
- (ii) **Coordination:** If $\kappa > \rho$, the best equilibrium outcome (x^κ, y^κ) is implementable. $\kappa \geq \rho$ is a necessary condition for implementation.
- (iii) Otherwise, when $\kappa < \rho$, the welfares in the best and the worst equilibrium outcomes, $\bar{w}(x^\kappa, y^\kappa)$ and $v_{\text{worst}}(\kappa) = \min v(\kappa)$, are weakly increasing in κ .

PROOF. Points (i) and (iii). Let us start by showing that the set of equilibrium payoffs is a compact set, that (x^κ, y^κ) is an equilibrium outcome that achieves the best equilibrium payoff and that both the worst and the best payoffs are weakly increasing with κ .

¹³Note that multiple $y^*(x)$ may solve the ex post government problem.

¹⁴That controllability implies $\rho < \infty$ results from X and Y being compact sets and w being a continuous function.

To this purpose, let us show the following lemma. Let

$$S_\kappa = \{(x, y) \in C, w(x, y^*(x)) - w(x, y) \leq \kappa\}$$

be the set of competitive outcomes that can be ex post sustained by the government if the reaction function is such that $\bar{y}(x) = y$. Let

$$S_\kappa^x = \{x \in X, \exists y \in D(x), (x, y) \in S_\kappa\}$$

be the set of private-sector actions that belong to at least one allocation in S_κ .

LEMMA 2. *An allocation (x, y) is an equilibrium outcome if and only if:*

$$(i) \quad (x, y) \in S_\kappa,$$

$$(ii) \quad \bar{w}(x, y) \geq \min_{x' \in S_\kappa^x} \max_{y' | (x', y') \in S_\kappa} \bar{w}(x', y').$$

PROOF. Suppose that (x, y) is an equilibrium outcome. This means that there exists $\sigma = (\bar{y}, \sigma^g, \sigma^h)$ such that σ leads to (x, y) . By definition, $(x, y) \in C$ and $w(x, y^*(x)) - w(x, y) \leq \kappa$. As a result, $(x, y) \in S_\kappa$.

Let us show that condition (ii) is also satisfied. Indeed, suppose it is not. Without loss of generality, we can focus on $\bar{y}(x) = y$. In this case, for all $x' \in S_\kappa^x$, there exists y' such that $(x', y') \in S_\kappa$, such that $\bar{w}(x', y') > \bar{w}(x, y)$. Using these values of x' and y' , we can build a reaction function $\bar{\eta}$ such that $\bar{\eta}(x') = y'$ on S_κ^x —outside this set, we can take $\bar{\eta}$ in an arbitrary manner as no equilibrium can form. The strategy profile σ should define the actions after $\bar{\eta}$: $\sigma^h(\bar{\eta}) = x'$ and $\sigma^g(\bar{\eta}, \sigma^g(\bar{\eta})) = y'$.

As $\bar{w}(x, y)$ is the payoff associated with σ , that is, $\bar{r}(\bar{y}, x, y)$ and the payoff associated with $\bar{\eta}$ is $\bar{r}(\bar{\eta}, x', y') = \bar{w}(x', y')$ (notice that we have constructed $\bar{\eta}$ so that $\bar{\eta}(x') = y'$), which implies that there exists a reaction function $\bar{\eta}$ that allows to do better than $\bar{y}$, a contradiction.

Conversely, suppose that (x, y) satisfies conditions (i) and (ii) from Lemma 2. Let us show that it is an equilibrium outcome. Let us consider the reaction function such that $\bar{y}(x) = y$ and the strategies so that $\sigma^h(\bar{y}) = x$ and $\sigma^g(\bar{y}, \sigma^h(\bar{y})) = y$. As $(x, y) \in S_\kappa$, $\sigma^h(\bar{y}) = x$ and $\sigma^g(\bar{y}, \sigma^h(\bar{y})) = y$ form a competitive outcome and y is optimal after $\sigma^h(\bar{y})$.

Let us consider $\bar{\eta} \in Y(X)$ such that $\mathcal{CE}(\bar{\eta})$ is nonempty. For any $x' \in S_\kappa^x$, by definition, $\bar{w}(x', \bar{\eta}(x')) \leq \max_{y' | (x', y') \in S_\kappa} \bar{w}(x', y')$. As $\sigma^h(\bar{\eta})$ can be chosen in S_κ^x and then $\sigma^g(\bar{\eta}, x')$ can be selected such that $(x', \sigma^g(\bar{\eta}, x')) \in S_\kappa$, with

$$\bar{r}(\bar{\eta}, \sigma^h(\bar{\eta}), \sigma^g(\bar{\eta}, \sigma^h(\bar{\eta}))) \leq \min_{x' \in S_\kappa^x} \max_{y' | (x', y') \in S_\kappa} \bar{w}(x', y').$$

Using condition (ii), the right-hand term can be bounded above by $\bar{w}(x, y)$, which is the payoff associated with $\bar{r}(\bar{y}, x, y)$, thus showing that there exists a strategy profile such that it is optimal to play $\bar{y}$. As a result, (x, y) is an equilibrium outcome. $\square$

From Lemma 2, v is a compact as S_κ is a compact and condition (ii) involves an inequality that is not strict. Furthermore, $(x^\kappa, y^\kappa) \in \Theta(\kappa)$: indeed, it belongs to S_κ and leads to the highest payoff in S_κ so that (ii) is trivially satisfied.

The worst equilibrium outcome Let us now consider the worst equilibrium outcome $v_{\text{worst}}(\kappa)$ for a given κ . The corresponding strategies are $y(\cdot)$, σ^g , and σ^h . For a larger κ , let us note that these strategies still satisfy condition (i) of the definition of an equilibrium but, given the larger commitment ability, the government may sustain something at least better. In particular, this means that something worse than $v_{\text{worst}}(\kappa)$ is not an equilibrium outcome. So, the worst equilibrium outcome is an increasing function of κ .

Point (ii) The following lemma shows under which conditions the best equilibrium outcome is implementable.

LEMMA 3. *The best equilibrium outcome (x^κ, y^κ) is implementable if and only if there exists a reaction function $\bar{y}$ such that:*

- (i) $\bar{y}(x^\kappa) = y^\kappa$,
- (ii) $\forall x' \neq x^\kappa, (x', \bar{y}(x')) \notin C$,
- (iii) $\forall (x', y') \in C, (x', y') \neq (x^\kappa, y^\kappa), w(x', y') - \kappa < w(x', \bar{y}(x'))$.

PROOF. First, suppose that there exists a reaction function $\bar{y}$ such that $\bar{y}(x^\kappa) = y^\kappa$ and

$$\forall x' \neq x^\kappa, (x', \bar{y}(x')) \notin C; \quad (10)$$

$$\forall (x', y') \in C, (x', y') \neq (x^\kappa, y^\kappa), w(x', y') - \kappa \leq w(x', \bar{y}(x')). \quad (11)$$

First, notice that committing to such a reaction function is part of an equilibrium that leads to an equilibrium outcome (x^κ, y^κ) . Furthermore, there does not exist any alternative equilibrium outcome. Indeed, suppose that such an equilibrium outcome exists $(x', y') \neq (x^\kappa, y^\kappa)$. This means that there exists a reaction function $\bar{\eta}$ such that $\bar{\eta}(x) = y$ and a strategy profile σ^h and σ^g such that $\sigma^h(\bar{\eta}) = x$ and $\sigma^g(\bar{\eta}, \sigma^h(\bar{\eta})) = y$ and $\bar{\eta}$ is played in equilibrium. However, this strategy profile should be also such that $\sigma^h(\bar{y}) = x^\kappa$ and $\sigma^g(\bar{y}, \sigma^h(\bar{y})) = y^\kappa$. Indeed, from Condition (11), the government is better off to stick to $\bar{y}$ when having committed to it and (x^κ, y^κ) is the only continuation of an equilibrium that can form after $\bar{y}$. Finally, as (x^κ, y^κ) leads to the highest payoff, the government strictly prefers to play $\bar{y}$ instead of $\bar{\eta}$.

Let now prove the reciprocal. Suppose that (x^κ, y^κ) is implementable. Let us show that there exists a reaction function satisfying the three conditions of Lemma 3.

As (x^κ, y^κ) is an equilibrium outcome, there exists $\bar{y}$ such that $(\bar{y}, \sigma^h, \sigma^g)$ is an equilibrium.

First, condition (i) is satisfied. Either $\bar{y}(x^\kappa) = y^\kappa$ or, if not, we can consider another $\bar{y}'$ that coincides with $\bar{y}$ except for $x = x^\kappa$, where $\bar{y}'(x^\kappa) = y^\kappa$ and this alternative reaction function yields a strictly higher payoff as the cost κ is not incurred in equilibrium.

Second, suppose that conditions (ii) and (iii) are not satisfied. Let us show there exists a competitive outcome (x, y) that the government cannot rule out with any reaction function. Indeed, if the reciprocal of Lemma 3 is not true, for all reaction functions $\bar{\eta}$, either there exists $x \neq x^\kappa$ such that $(x, \bar{\eta}(x)) \in C$ or there exists y such that (x, y) is such that $w(x, y) - \kappa > w(x, \bar{\eta}(x))$, where $(x, y) \in C$.

The latter case is possible if and only if the ex post optimal action $y^*(x)$ satisfies $w(x, y^*(x)) - \kappa > w(x, \bar{\eta}(x))$. Otherwise, the government can select the reaction function to be $y^*(x)$ to rule out x . As a result, we can consider the strategy for the private sector to play $\sigma^h(\bar{\eta}) = x(\bar{\eta}) \neq x^\kappa$ and for the government to play $\sigma^g(\bar{\eta}, \cdot) = y^*(\cdot)$. This is so that the continuation of an equilibrium after the commitment to a reaction function intended to lead to (x^κ, y^κ) is $(x(\bar{\eta}), y(\bar{\eta}))$ with $y(\bar{\eta}) = \bar{\eta}(x(\bar{\eta}))$ if $(x(\bar{\eta}), \bar{\eta}(x(\bar{\eta}))) \in C$ or $y(\bar{\eta}) = y^*(x(\bar{\eta}))$ otherwise.

Finally, let us note that, anticipating a private sector strategy that would lead to an inferior outcome than (x^κ, y^κ) , the government may simply commit to another reaction function but that cannot lead for sure to (x^κ, y^κ) , so that another equilibrium outcome exists, contradicting implementation. $\square$

First, a simple application of Lemma 3 shows that if the best equilibrium outcome is implementable for some $\underline{\kappa} > 0$, then it is also implementable for any $\kappa \geq \underline{\kappa}$.

Second, Lemma 3 shows that when $\kappa > \rho$ the best equilibrium outcome is implementable. Construct $\bar{y} \in Y(X)$ that satisfies the three conditions. Define $\bar{y}$ such that $\bar{y}(x^\kappa) = y^\kappa$. Item (i) is satisfied. Lemma 3 and the definition of ρ mean that, for any $x \in A_X$, we can find $\bar{y}(x) \in \mathcal{Y}(x)$ such that $w(x, y^*(x)) - \rho \leq w(x, \bar{y}(x))$. Hence, $(x, \bar{y}(x)) \notin C$ which leads to item (ii). Since $y^*(x)$ is a best response to x , for any couple $y \in D(x)$ such that $(x, y) \in C$, $w(x, y) - \rho \leq w(x, \bar{y}(x))$. Then $\kappa > \rho$ implies item (iii). Therefore, the best equilibrium outcome is implementable.

Reciprocally, suppose that the best equilibrium outcome is implementable. Thus, for any $x \neq x^\kappa \in X$, we can build a $\bar{y}(x) \in \mathcal{Y}(x)$ and $w(x, y) - \kappa < w(x, \bar{y}(x))$ for any $(x, y) \in C$. Thus, it is also true for $y = y^*(x)$, and hence, $w(x, y^*(x)) - \kappa < w(x, \bar{y}(x))$. This proves that $\kappa \geq \rho$. $\square$

As a result of Proposition 1, the main variable to examine in determining the extent to which a government can solve a coordination problem is ρ : a unique equilibrium obtains when $\kappa > \rho$, and this equilibrium is the best time-consistent competitive outcome (x^κ, y^κ) defined above.

The sketch of the proof detailed above goes as follows. On the one hand, the government can commit to a reaction function such that the response to x^κ is y^κ , that is, $\bar{y}(x^\kappa) = y^\kappa$. Such a commitment implies that the deviation from y^κ would lead to a cost κ for the government, thus making y^κ a time-consistent action for the government. On the other hand, when $\rho < \kappa$, the government commits to play $\bar{y}(x) \in \mathcal{Y}(x)$, which discourages agents from playing x whenever $x \neq x^\kappa$, as in [Bassetto \(2005\)](#). $\rho < \kappa$ ensures that $\bar{y}(x)$ can be chosen so that the government is ex post better off playing that action rather than anything else.

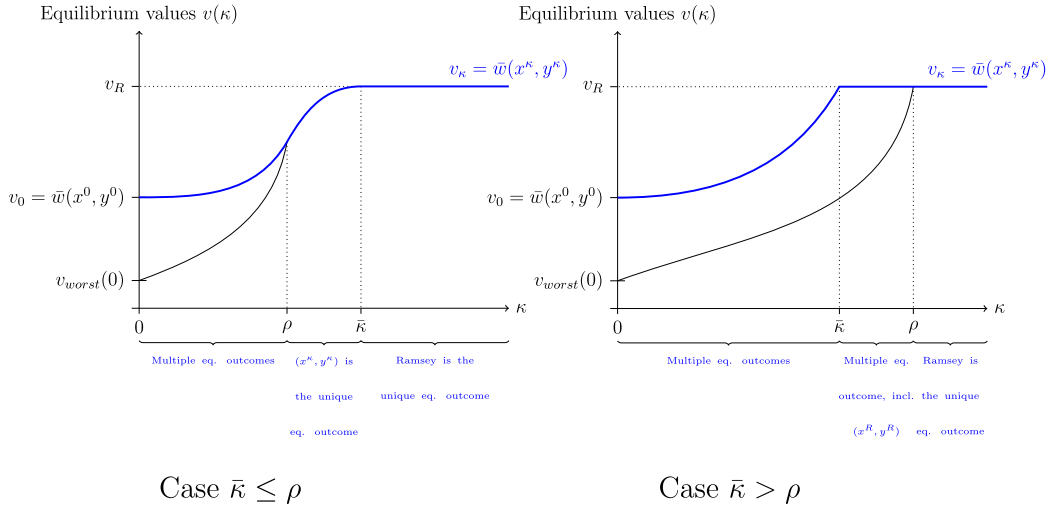


FIGURE 2. Best and worst equilibrium payoffs as a function of κ . In the two graphs, we plot the best (v_κ) and the worst (v_{worst}) equilibrium payoffs for the government as a function of commitment ability (κ). The left-hand panel corresponds to the case in which $\bar{\kappa}$ —the minimum commitment ability required to achieve the Ramsey outcome—is below the cost of controllability ρ . The right-hand panel corresponds to the case in which $\bar{\kappa}$ is above ρ .

The proposition also clarifies that the best time-consistent competitive outcome is always an equilibrium outcome; one that leads to the best equilibrium payoff for the government.

Finally, the proposition provides some comparative statics on the equilibrium set as a function of commitment ability. Figure 2 summarizes the result of Proposition 1 in the two cases in which $\rho \leq \bar{\kappa}$ and $\rho > \bar{\kappa}$, where $\bar{\kappa}$ is the minimum level of commitment ability such that the Ramsey allocation is an equilibrium outcome and $(x^\kappa, y^\kappa) = (x^R, y^R)$. As the figure illustrates, increasing κ improves the equilibrium outcomes, not only the best one but also the worst one.

Full commitment ability A direct implication of Proposition 1 is when the government can fully commit, that is, when $\kappa = \infty$.

COROLLARY 4 (Bassetto (2005), Atkeson, Chari, and Kehoe (2010)). *Suppose that $\kappa = \infty$ and Assumption 2 is satisfied. Then the government implements the Ramsey outcome, $\Theta(\infty) = \{(x^R, y^R)\}$.*

When $\kappa = \infty$, κ is larger than $\bar{\kappa}$. The Ramsey allocation is an equilibrium outcome and, by definition, the best one. When $\kappa = \infty$ and Assumption 2 is satisfied, the cost of controllability is well-defined and κ is larger than ρ ; the government can credibly rule out any inferior outcomes.

How much additional commitment is needed to solve the coordination problem? The critical variable for the ability of the government to solve the coordination problem is ρ .

This variable captures the welfare cost for the government to engage in actions that discourage individual private agents from playing something other than the government desires. In the following proposition, we characterize this cost of controllability depending on whether the action set is continuous or discrete.

PROPOSITION 5. *The cost of controllability satisfies:*

- (i) **Continuous action set:** *If any Nash outcome $(x, y) \in \Theta(0)$ is interior in $A_X \times A_Y$ and $u_{13}(x, x, y) \neq 0$ for any Nash outcome, then $\rho = 0$.*
- (ii) **Discrete action set:** *When the action set is a regular grid with steps Δ_x and Δ_y for x and y , if any Nash outcome $(x, y) \in \Theta(0)$ is sufficiently interior and*

$$\kappa_{11} > |u_{11}|, \quad \text{and} \quad |u_{13}| > \kappa_{13} > 0,$$

$$\text{then } \rho \leq \max |w_2| \max[2 \frac{\kappa_{11}}{\kappa_{13}} \Delta_x, \Delta_y].$$

PROOF. We start this proof by showing that it is necessary and sufficient to only deter Nash outcomes.

LEMMA 6. *The government implements (x^κ, y^κ) if and only if it can credibly deter the private sector from playing the actions $x \in A_X$ different from x^κ and consistent with a Nash outcome. More formally,*

$$\rho = \max_{x|(x,y) \in \Theta(0)} \min_{y \in \mathcal{Y}(x)} [w(x, y^*(x)) - w(x, y)],$$

where $y^*(x) \in \arg \max_{y \in D(y)} w(x, y)$.

PROOF. Lemma 6 directly results from the definition of ρ . Take $x \neq x^\kappa$ that is not consistent with a Nash outcome, which means, such that $y^*(x) \notin C$. Therefore, the cost of deterring the private sector from playing this action x is zero as $y^*(x) \in \mathcal{Y}(x)$. So, if the government can credibly deter the private agents from playing actions $x \neq x^\kappa$ that are consistent with a Nash outcome, the government implements (x^κ, y^κ) . This condition means that for any $x \neq x^\kappa$ there exists a policy action y in $\mathcal{Y}(x)$ such that

$$w(x, y) > w(x, y^*(x)) - \kappa. \quad (12)$$

□

Regarding the continuous case, we show that the proposition ensures the existence of a reaction function satisfying the conditions of Lemma 3, which is sufficient to guarantee implementation. For any $x \neq x^\kappa$, such that there exists $y^*(x) \in \arg \max_{y \in D(y)} w(x, y)$ so that $(x, y^*(x)) \notin C$, the reaction function will be simply the best response $y^*(x)$ as explained in Lemma 6. In the other cases, all the actions $x \neq x^\kappa$ are such that, for any $y^*(x) \in \arg \max_{y \in D(y)} w(x, y)$, $(x, y^*(x)) \in C$. In these cases, the government cannot count on one of its ex post best responses to deter private agents from

playing $\xi = x$. We assume that all the actions of these Nash outcomes are interior actions. Take one best response $y^*(x)$. Since we assume that there exists an open interval around the best response $y^*(x)$ in $D(x)$, we can perturb the best response in $D(x)$ by a small amount $\epsilon(x) \geq 0$ such that

$$\exists \xi \in A_X, \quad u(\xi, x, y^*(x) + \epsilon(x)) > u(x, x, y^*(x) + \epsilon(x)), \quad (13)$$

$$\text{and} \quad |w(\tilde{x}, y^*(x) + \epsilon(x)) - w(x, y^*(x))| \leq \kappa/2. \quad (14)$$

Notice that, for any allocation $(x, y) \in \Theta(0) \setminus \{(x^\kappa, y^\kappa)\}$, because u is twice continuously differentiable, strictly concave in ξ , and $u_{13}(x, x, y) \neq 0$ the implicit function theorem applied to $u_1(\xi^*, x, y)$ allows to write, given that $u_{11} < 0$,

$$\left. \frac{d\xi^*}{dy} \right|_{\xi^*=x} = -\frac{u_{13}(x, x, y)}{u_{11}(x, x, y)} \neq 0,$$

thus, showing that a marginal change of y induces a change of ξ^* .

As a result, that $d\xi^*/dy|_{\xi^*=x} \neq 0$ and x in the interior of A_X means that we can pick $\epsilon(x)$ satisfying the first inequality that is arbitrarily close to 0 and so that (i) $y^*(x) + \epsilon(x) \in A_Y$ and (ii) the second inequality is satisfied simultaneously by continuity of w . Now construct the reaction function $\bar{y} \in Y(X)$ as follows:

$$\begin{aligned} \bar{y}(x) &= y^\kappa, & \text{if } x = x^\kappa \\ \bar{y}(x) &= y^*(x) + \epsilon(x), & \text{otherwise,} \end{aligned}$$

where $\epsilon(x) = 0$ when $(x, y^*(x)) \notin C$ or equivalently $y^*(x) \in \mathcal{Y}(x)$. Equation (13) means that for any $x' \neq x^\kappa$, $(x', \bar{y}(x')) \notin C$; Equation (14) ensures that the third item of Lemma 3 is verified. Therefore, $\bar{y}$ satisfies the three conditions of Lemma 3. The best equilibrium outcome is thus implementable.

Let us finally prove the Proposition 5 in the discrete case. We define the action set as follows: $A_X = \{x_1, \dots, x_N\}$ and $A_Y = \{y_1, \dots, y_P\}$ where N and P are some large positive integers. We proceed as for the continuous action set and construct a reaction function that coincides with the above reaction function everywhere except for the aggregate action x such that $y^*(x) \notin \mathcal{Y}(x)$ and x is part of a Nash outcome (i.e., there exists $y \in D(x)$ such that $(x, y) \in \Theta(0) \setminus \{(x^\kappa, y^\kappa)\}$). Take such a private-sector action $x = x_n$. We denote one of the best responses $y^*(x) = y_{n^*}$. What we show is that there exists at least an action y_{n^*+p} with some p such that

$$u(x_{n+1}, x_n, y_{n^*+p}) > u(x_n, x_n, y_{n^*+p}).$$

As we assume that the Nash outcome is sufficiently interior, $x_{n+1} \in A_X$. We prove this claim by assuming that u_{13} is positive, but the same logic applies if it is negative. By application of the mean value theorem, there exists some $x' \in (x_n, x_{n+1})$, such that

$$u(x_{n+1}, x_n, y_{n^*+p}) - u(x_n, x_n, y_{n^*+p}) = u_1(x', x_n, y_{n^*+p})\Delta x.$$

A second application of the same theorem and the fact that $u_{13} > \kappa_{13}$, leads to

$$u(x_{n+1}, x_n, y_{n^*+p}) - u(x_n, x_n, y_{n^*+p}) > [u_1(x', x_n, y_{n^*}) + \kappa_{13} p \Delta_y] \Delta_x.$$

Notice also that because we have $u(x_n, x_n, y_{n^*}) > u(x_{n-1}, x_n, y_{n^*})$ and $u(x_n, x_n, y_{n^*}) > u(x_{n+1}, x_n, y_{n^*})$, it must be the case that there exists a maximum in (x_{n-1}, x_{n+1}) . As a consequence (using always the mean value theorem), we get $2\Delta_x \kappa_{11} > |u_1(x', x_n, y_{n^*})|$. Therefore, x_{n+1} is preferred to x_n by private agents when

$$p \Delta_y \geq \max \left[\frac{2\kappa_{11}}{\kappa_{13}} \Delta_x, \Delta_y \right]. \quad (15)$$

The second term results from the fact that p is an integer that cannot be below: 1. Therefore, if the government commits to play y_{n^*+p} after x_n , private agent will not find optimal to play $\xi = x_n$ but will instead play $\xi = x_{n+1}$. As we assume that the Nash outcome is sufficiently interior, $y_{n^*+p} \in A_Y$. It remains to compute the associated ex post cost for the government of such an action. As $w(\cdot, \cdot)$ is strictly concave and twice differentiable, we get the following upper bound:

$$w(x_n, y_{n^*}) - w(x_n, y_{n^*+p}) \leq |w_2(x_n, y_{n^*})| p \Delta_y \leq \sup |w_2| \max \left[\frac{2\kappa_{11}}{\kappa_{13}} \Delta_x, \Delta_y \right].$$

Notice that when u_{13} is negative the marginal cost to consider in the above inequality is $|w_2(x_n, y_{n^*-p})|$, otherwise the logics remains the same, which leads to the upper bound for ρ in the Proposition 5. The $\max |w_2|$ in the proposition is then the maximum of the derivate of w with respect to y over all the grid points for y and for the admissible $x \in A_X \setminus \{x_\kappa\}$ belonging to a Nash outcome such that $y^*(x) \notin \mathcal{Y}(x)$. $\square$

Proposition 5 relies on the fact that only the private-sector actions belonging to a Nash outcome may be costly to deter, that is, contribute to the cost of controllability ρ . This is the first step of the proof (Lemma 6).

Proposition then considers two situations: a continuous set of actions and a discrete set of actions. In both cases, we focus on interior Nash outcomes. In the continuous case, formally, a Nash outcome $(x, y) \in \Theta(0)$ is interior when there exists an open interval $I(x) \subset D(x)$ containing y and when $x \in \overset{\circ}{X}$, the interior of X . In the discrete case, a stronger condition is required: $x - \Delta_x$ and $x + \Delta_x$ have to be in the action set A_x . In addition, if we denote by n^* the index of the best response $y^*(x) \in A_Y$, we require that y_{n^*+p} (or depending on the sign of u_{13} , y_{n^*-p}) is in A_Y with p the smallest integer such that

$$p \Delta_y \geq \max \left[\frac{2\kappa_{11}}{\kappa_{13}} \Delta_x, \Delta_y \right].$$

Continuous action sets When the action set is continuous, Proposition 5 shows that, if all the Nash outcomes are *interior* and if the marginal utility of private agents locally depends on the policy action, then the government can deter all the actions associated with a Nash outcome, and hence, the government can solve the coordination problem at no cost.

To see that, take an interior Nash outcome (x, y) such that $u_{13}(x, x, y) > 0$. As x is interior, the marginal utility of an individual agent is zero for the allocation $(u_1(x, x, y) = 0)$, as otherwise, this agent would be better off by setting $\xi \neq x$. Given that $u_{13}(x, x, y) > 0$, the government can select an action slightly above the ex post best action y and increase the individual agent's marginal utility above 0; that is, $u_1(x, x, y) > 0$. In turn, the individual agent is better off increasing its action ξ above x to maximize utility, which is feasible since the action x belongs to the interior of X . Finally, the policy action is almost not costly since it can be selected arbitrarily close to the ex post best action. The same reasoning applies when $u_{13}(x, x, y) < 0$, and more generally, when $d\xi^*/dy|_{\xi^*=x} \neq 0$.

Using these elements, we can build a reaction function $\bar{y}$ to which the government can credibly stick and so that only the best competitive outcome can form (x^κ, y^κ) . Notice that the reaction function that we build in the proof of Proposition 5 may have to be discontinuous with respect to the private sector action x . However, discontinuity is not always required to implement the best competitive outcome and continuous reaction function can be sometimes be designed. As we discuss below in examples, the reaction function so that almost no commitment ability is needed is discontinuous in the capital taxation example but can be continuous in the bailout example. We discuss further these results and how, for example, imperfect information requires to focus on continuous reaction functions in the working paper version of the paper (Barthélemy and Mengus (2022)).

When a Nash outcome is not interior, the cost of controllability may be high ($\rho > 0$) even if $d\xi^*/dy|_{\xi^*=x} \neq 0$. The government may be unable to depart from its ex post best action locally in the right direction (in the above paragraph, choosing an action y slightly above the ex post optimal one $y^*(x)$) because of a feasibility constraint on its action or because the private sector cannot move away from the aggregate private action x locally in the right direction (in the above paragraph, choosing ξ slightly above x). In such a case, and under the controllability assumption, only a costly policy action may succeed in deterring private agents from playing the undesired action $x \neq x^\kappa$, leading to a high cost of controllability.

Discrete action sets When the action set is discrete, the above reasoning does not apply as the government cannot count on a marginal change to induce a marginal change from private agent. So, the cost of controllability ultimately depends on the distance between two actions for the government (Δ_y) and for private agents (Δ_x). Not very surprisingly, when these two distances tend to 0, the second item of the proposition shows that $\rho = 0$ coinciding with the first item of the proposition. Otherwise, the cost of controllability is small when private agents' marginal utility is very sensitive to government's action (u_{13} large), when they change a lot their actions because of a change in marginal utility (u_{11} small) and when government's actions are sufficiently precise (Δ_y small). Notice also that contrary to the continuous case, the objective function w matters as the government has to commit and stick to an action that may be far away from its ex post best response involving an ex post cost that depends on the general slope of w with respect to its action.

In the next section, we illustrate how Proposition 5 applies in different examples, such as the Farhi and Tirole (2012) model of bailouts and the capital taxation problem.

4. EXAMPLES

The model laid out above can encompass multiple macroeconomic situations under discretion ($\kappa = 0$). In this section, we describe some of these situations and illustrate how time inconsistency leads to coordination problems. To start with, the model of bailouts by [Farhi and Tirole \(2012\)](#) illustrates how time inconsistency leads to a coordination problem. This happens even though the Ramsey outcome is an equilibrium outcome. Second, under some conditions, a simple model of capital taxation, as in [Chari and Kehoe \(1990\)](#), illustrates a situation of time inconsistency leading to a coordination problem where the Ramsey outcome is not an equilibrium outcome. Finally, we show how, in these two examples, the government can obtain a unique equilibrium outcome with an arbitrarily small commitment ability.

4.1 Bailout problem

The environment Consider a bailout problem as in [Farhi and Tirole \(2012\)](#), from which we borrow the notation. There are three periods $t \in \{0, 1, 2\}$. The economy is populated by risk-neutral bankers, deep-pocket risk-neutral investors, and a government.

At date 0, bankers receive an endowment A , and they can invest in a risky investment opportunity and borrow short-term. At date 1, with probability α , the investment opportunity yields $(\pi + \rho_1)i$ for i invested at date 0, among which $(\pi + \rho_0)i$ can be pledged to investors with $\rho_0 < 1$. With probability $1 - \alpha$, investment yields only πi at date 1 and $\rho_1 j$ at date 2 with $j \leq i$ the amount of resources reinvested at date 1. In this case, only $\rho_0 j$ can be pledged to investors. The returns on investment opportunities are perfectly correlated across bankers. At date 0, bankers optimally set a contingent short-term debt contract equal to πi in the absence of crisis and di otherwise.

The government sets the real rate of interest. Between date 0 and date 1, as well as between date 1 and date 2 when investment is successful, the government optimally selects an interest rate equal to 1. In the event of a crisis, the government sets an interest rate $R \leq 1$ between date 1 and date 2 to maximize its objective function:

$$-L(R) - \frac{(1-R)\rho_0 j}{R} + \beta j \quad (16)$$

with $L(R)$ a deadweight loss associated with setting interest rates below 1. L satisfies $L(R) \geq 0$, $L(1) = L'(1) = 0$ and L is decreasing on $[\rho_0, 1]$. The second term of (16) corresponds to the subsidy from savers to borrowing banks at a rate below 1. The last term stands for the gain due to higher date-1 reinvestment—by convention, $j = i$ in the case of a successful investment. The date-0 objective of the government is the expectation of the date-1 objective function. In this model, bankers play first by selecting investment i and short-term debt d . Then the government plays R at date 1 (in the event of a crisis), and then the bankers decide to reinvest if needed.

When investment is unsuccessful and needs reinvestment, bankers optimally select reinvestment j so that

$$j = \min \left\{ \frac{\pi - d}{1 - \frac{\rho_0}{R}}; 1 \right\} i. \quad (17)$$

At date 0, this leads bankers to select investment i so that

$$i = \frac{A}{1 - \pi - \alpha\rho_0 + (1 - \alpha)\xi}, \quad (18)$$

where $\xi \equiv \pi - d$ is the liquidity ratio; that is, ξi is the banker's cash-flow available at date 1 in case of a crisis net of debt repayment di . ξ maximizes

$$(\rho_1 - \rho_0)(\alpha i + (1 - \alpha)j) = (\rho_1 - \rho_0) \left(\frac{\alpha + (1 - \alpha) \min \left\{ \frac{\xi}{1 - \rho_0/R}; 1 \right\}}{1 - \pi - \alpha\rho_0 + (1 - \alpha)\xi} \right) A. \quad (19)$$

When $\alpha + \pi < 1$ as assumed by Farhi and Tirole (2012), the maximization of (19) leads to $x = 1 - \rho_0/R$.

Mapping with the general model Farhi and Tirole's (2012) model is a game between bankers at date 0 and the government's intervention in the case of unsuccessful investment at date 1. Notice that the individual banker's action ξ does not depend on other actions, except through the dependence of the policy rate R on the aggregate bankers' decisions. Using ξ as defined above, we can then map this model to our general setting as follows:

$$\xi \in [0, 1 - \rho_0], \quad x \in [0, 1 - \rho_0], \quad y = R, \quad \text{and} \quad D(x) = [\rho_0, 1].$$

We can express the objective function (16) as a function of $y = R$ and $x = \xi$ by using (17) and (18). This defines $w(x, y)$. Finally, (19) defines $u(\xi, x, y)$.

In this example, the objective function relevant ex ante for the government differs from the one relevant ex post due to the uncertainty around the success of the investment at date 1. Ex ante, this objective function is

$$\bar{w}(x, y) = \alpha\beta i(x) + (1 - \alpha) \left(-L(y) - \frac{(1 - y)\rho_0 j(x)}{y} + \beta j(x) \right).$$

First, under the assumption that β is sufficiently small ($\beta \leq 2 - \alpha - \pi - \rho_0$), the Ramsey allocation is such that $R = 1$ (Proposition 1 in Farhi and Tirole (2012)). The Ramsey allocation is an equilibrium outcome in this model. This happens under the condition that $(\beta + \rho_0 - 1)A/(1 - \pi - \alpha\rho_0) \geq L(\rho_0)$ (Corollary 1 in Farhi and Tirole (2012)). Second, there are multiple other equilibria with bailout, $R < 1$ (see Proposition 2 in their paper): this model features a coordination problem, even though time inconsistency does not prevent the Ramsey allocation from being an equilibrium outcome.

The equilibrium set under limited commitment ability Our general results imply that, in this model, for any positive—and potentially arbitrarily low— $\kappa > 0$, the Ramsey allocation is the unique equilibrium outcome. Indeed, the inspection of the optimal decision x indicates that it is such that $d\xi^*/dy|_{\xi^*=x} \neq 0$ and that the Nash outcomes, with the exception of the Ramsey outcome, are interior. As a result, the conditions of Proposition 5 are satisfied.

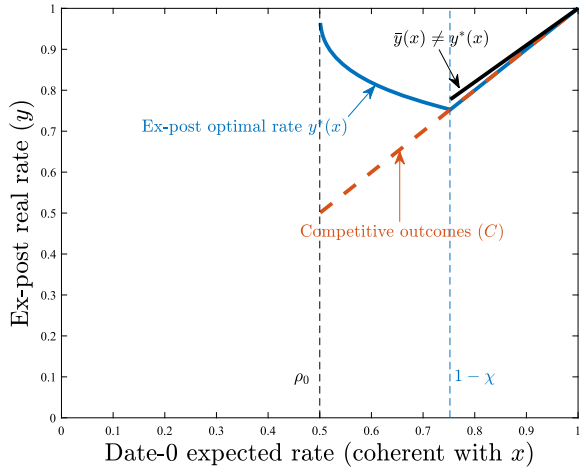


FIGURE 3. Obtaining a unique equilibrium in Farhi and Tirole (2012) model of bailouts. In this graph, we plot the set of competitive outcomes (C), the ex post optimal response by the government as a function of private-sector actions x ($y^*(x)$), and finally, one reaction function $\bar{y}$ that allows the government to select a unique equilibrium outcome when this reaction function differs from the ex post optimal reaction function y^* . Notice that such a reaction function $\bar{y}$ allows the government to implement the Ramsey allocation.

But how does this work in practice and how can the government deter bankers from anticipating a bailout? Figure 3 illustrates what the government may do.

In this model, the ex post optimal policy is to set the interest rate R at the level at which the private sector expected it, at least when this interest rate does not deviate much from the Ramsey level $R = 1$. The cost of decreasing R is continuous in R , while there is a fixed gain related to refinancing projects. This leads to a continuum of equilibria. To prevent this, the government simply commits to a bailout policy very close to the one expected— $R = R^e + \epsilon$ —where R^e is the expected rate by the private sector. Then the government calibrates ϵ as a function of its commitment ability κ : $w(x(R^e), R + \epsilon) \leq w(x(R^e), R) + \kappa$. As can be observed in Figure 3, such a bailout policy prevents any equilibrium to form where $R < 1$. Notice that nothing prevents to also restrict to continuous reaction functions with respect to the expected interest rate R^e .

How realistic are such partial bailouts? Implementing this solution may be quite involved; first, because it may require eliciting the private sector's expectations, and second, because it also requires that private agents are sufficiently responsive to small variations in the expected policy. However, in other contexts, such as the rescue of a sovereign like Greece during the euro area sovereign debt crisis, a form of partial bailouts was put in place with private-sector involvement (PSI). Our result suggests that committing to a limited private-sector involvement may be sufficient to rule out expectations of bailout in equilibrium. In addition, such a low-cost commitment to partial bailouts make nonnecessary Farhi and Tirole's (2012) solution, which requires regulating bankers by imposing a cap on short-term debt, or equivalently, a liquidity requirement at date-0, whenever the government cannot perfectly commit.

4.2 Optimal capital taxation problem

The environment Consider a two-period taxation problem (adapted from Fischer (1980), Chari and Kehoe (1990), Bassetto (2005)). Time is discrete and indexed by $t = \{1, 2\}$. The economy is populated by a continuum of households and a government. At date 1, each household receives an endowment ω and decides to consume $\tilde{c}_1$ or to invest ξ in a linear saving technology, which yields $R(k)\xi > 1$ units of goods at date 2, where $R(k)$ is a decreasing function of aggregate capital k . At date 2, households work and we denote by l the corresponding number of hours. We assume that the marginal product of labor is 1. The government can tax the return of capital at a tax rate δ and labor income at a rate τ . Then each household consumes the after-tax return of its investment and labor. Households value the consumption profile $(\tilde{c}_1, \tilde{c}_2)$ and labor $\tilde{l}$ using the utility function $\tilde{u}(\tilde{c}_1, \tilde{c}_2, \tilde{l})$. The government decides on taxes so as to maximize households' utility subject to the constraint to finance an exogenous amount of public expenditures G —its budget constraint is $G \geq \delta Rk + \tau l$, with l the aggregate number of hours.

The timing is as follows. First, the government commits to a reaction function $\bar{\delta}$, which maps a date-1 aggregate saving k to a (promised) tax rate $\delta = \bar{\delta}(k)$. Then households consume (in aggregate) c_1 and invest k . The government selects the tax on capital income δ and the tax on labor τ . If the government sets a tax rate different from $\delta = \bar{\delta}(k)$, it incurs a reneging cost κ . Finally, the households choose (in aggregate) l and consume c_2 .

Households select date-1 saving ξ for a given aggregate capital k and expected tax rates (δ, τ) as follows:

$$\begin{aligned} & \max_{\xi, l, c_1, c_2} \tilde{u}(c_1, c_2, l) \\ & \text{such that } c_1 \leq \omega - \xi \quad \text{and} \quad c_2 \leq R(k)\xi(1 - \delta) + l(1 - \tau). \end{aligned}$$

A competitive outcome must be such that $\xi = k$.

Under discretion ($\kappa = 0$), at date 2, that is, given k and c_1 , the government selects tax rates (δ, τ) by maximizing

$$\begin{aligned} & \max_{\delta, \tau, c_2, l} \tilde{u}(c_1, c_2, l) \\ & \text{such that } -\frac{u_l}{u_{c_2}} = 1 - \tau \quad \text{and} \quad G \leq \delta R(k)k + \tau l. \end{aligned}$$

The first constraint corresponds to the optimal consumption-leisure household's decision, the second to the government budget constraint. These two constraints implicitly define the labor tax $\tau(\delta, k)$ and the individual labor decision $l(\delta, k)$ as a function of the capital tax rate δ and the aggregate date-1 saving k . Notice that these functions are unaffected by the presence of a reneging cost. The date-2 decision by the government to tax labor and the households' decisions to consume and to work are nonstrategic. The strategic interaction concerns date-1 private saving decisions and date-2 capital taxation.

Mapping with the general model The model can be linked to the general setting defined above as follows:

$$\xi \in [0; \omega], \quad x = k \in [0; \omega], \quad y = \delta \in [0, 1], \quad D(x) = [0; 1],$$

$$\text{and } u(\xi, x, y) = \tilde{u}(\omega - \xi, R\xi(1 - y) + (1 - \tau(y, x))l(y, x), l(y, x)).$$

From the date-2 perspective, taxing capital is not distortive, while taxing labor is. Therefore, under discretion, the government taxes capital as much as needed to finance government expenditures. Under discretion, there exists an equilibrium in which households expect a tax rate $\delta = 1$ and do not save. There exists at least another equilibrium in which taxes on capital do not prevent households from saving when government's expenditures are low enough and the return on capital is high enough.

In the end, time inconsistency in the capital taxation model both prevents the Ramsey allocation from being an equilibrium outcome and leads to a coordination problem as multiple equilibria emerge.

The equilibrium set under limited commitment ability In this example, our general results help to rule out only interior Nash outcomes that are inferior for the government and the Nash outcome where capital is taxed at 100% is not an interior Nash outcome. Despite this, an arbitrarily small commitment ability is sufficient to rule out all inferior Nash outcomes.

Let us start with our general results. The controllability assumption depends on the combination of second derivatives of $\tilde{u}$ as follows:

$$u_{13}(x, x, y) = \frac{\partial l}{\partial y}(R(1 - y)\tilde{u}_{22} - \tilde{u}_{12} + R(1 - y)\tilde{u}_{23} - \tilde{u}_{13}) - R\tilde{u}_2. \quad (20)$$

When the utility function $\tilde{u}$ is separable in each argument, u_{13} is of the sign of $\partial l / \partial y R(1 - y)\tilde{u}_{22} - R\tilde{u}_2$, which is strictly negative when $\tilde{u}$ is strictly increasing and strictly concave in c_2 : an increase in the tax rate y decreases saving ξ . Note also that, in this case, u is also strictly concave in ξ . Under these assumptions, all interior Nash outcomes can be ruled out with an arbitrarily small commitment ability but this is not the case of the outcome where capital is fully taxed.

Let us illustrate how *all* the inferior Nash outcomes can be ruled out, including the one where capital is fully taxed. To this purpose, we calibrate the model as follows: $u(c_1, c_2, l) = 1/4 \log(c_1) + \beta(c_2 - l^2/2)$. Parameters are set to $R = \beta = 1$, $\omega = 1$, and $G = 0.1$.

Figure 4 plots the set of competitive outcomes (in red) and the ex post optimal policy (in blue)—the set of Nash outcomes is the intersection of the two. The best outcome is the one where capital is less taxed and savings are high. In black, we plot the reaction function that allows the government to rule out the two inferior outcomes. As we already discussed, this reaction function needs to differ from the ex post optimal policy only for the inferior outcomes. In this case, the government commits to a capital tax rate lower than the ex post optimal rate.

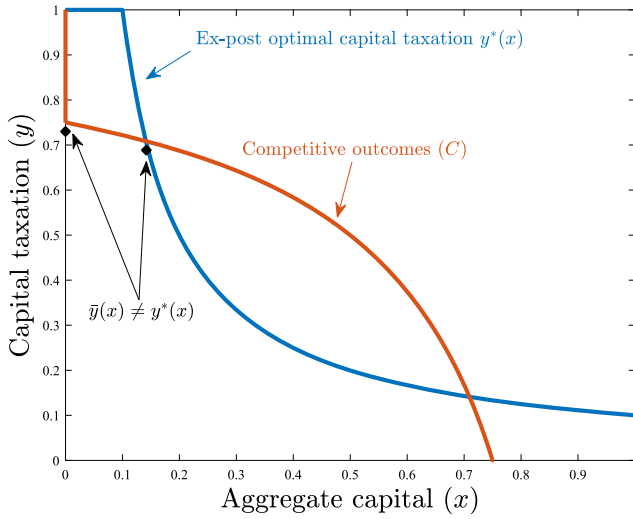


FIGURE 4. Equilibrium capital stock and capital taxation. This graph plots the set of competitive outcome (C) and the ex post optimal action $y^*(x)$ as a function of private agents' action x . In addition, we add to this graph the reaction function $\bar{y}$ that is sufficient to rule out all equilibria except the best one, when this reaction function differs from the ex post optimal reaction function y^* .

As can be observed, the interior Nash outcome can be ruled out by committing to an arbitrarily close tax rate, which makes the ex post cost of such a commitment arbitrarily small—thus requiring only an arbitrarily small commitment ability. However, the same logic cannot be applied to the Nash outcome where capital is fully taxed: the tax rate should be reduced by a substantial amount to push individual households to save. Yet, such an important reduction in the tax rate is not costly for the government: when aggregate capital is 0, modifying the tax rate yields 0 variations in the capital tax income received by the government.

In the end, the reaction function that we use in Figure 4 is discontinuous in two points: for $x = 0$ as we discussed above, but also close to the interior suboptimal Nash outcome. With such a discontinuous reaction function, the government does not require more than an arbitrarily small commitment ability to rule out equilibrium multiplicity. However, a natural question is whether such a result carries over when we restrict reaction functions to be continuous, for example, because governments may observe private actions only with a noise. We investigate this question formally in the working paper version of the paper (Barthélemy and Mengus (2022)). However, the main intuition can be grasped on Figure 4; one needs to draw on this graph a continuous reaction function that, on the one hand, crosses the set of competitive outcomes only in (x^κ, y^κ) and, on the other hand, is as close as possible from the ex post optimum action $y^*(x)$. Clearly, it is difficult not to draw this function to be below the set of competitive outcomes for $x < x^\kappa$, and so this function has to remain below these outcomes until $x = 0$

in order not to cross again the set of competitive outcomes. As a result, this continuous reaction function requires a strong commitment ability, as for low aggregate capital levels, it requires to play an action far away from the ex post optimal action $y^*(x)$.

5. DISCUSSION

In this section, we first discuss potential limits to our implementation results. We then provide interpretations of the cost of deviating from the reaction function, which allows the government to commit.

5.1 *Limits to the implementation result*

In this subsection, we describe extensions of the model that would imply a positive cost of controllability, $\rho > 0$, in contrast with the result of Proposition 5. These extensions are formally described in the working paper version of the paper (Barthélemy and Mengus (2022)).

Continuous reaction functions To prove the result of Proposition 5, we rely on discontinuous reaction functions. This approach proves to be critical in the capital taxation example in which, as we discussed, using continuous reaction functions only would require a strictly positive commitment ability. More generally, there can be good reasons to rely on continuous reaction functions. This is the case, for example, when the government can observe private actions only with some noise, being it arbitrarily small. Interestingly, imperfect information and continuous reaction functions do not always overturn our benchmark result on implementation, which holds, for example, in the bailout example whatever the level of noise. Therefore, the robustness of our benchmark results in the case of a continuous action set to the introduction of noises and to the use of continuous reaction functions depends on the precise details of the model and the exact topology of Nash equilibria. We keep such analysis for further research.

Fixed costs In the two economic models, apart from the reneging cost, payoffs are continuous. A fixed cost in private agents' maximization problem may lead to a positive cost of controllability, as it is the case with discrete action sets. The more expensive it is for private agents to pay the fixed cost, the more commitment ability the government needs to implement a unique outcome. On the contrary, as with discrete action sets, when the fixed cost is small enough, then the cost of controllability tends to zero under the assumptions of Proposition 5.

Repeated games In repeated settings, due to history-dependent strategies, the private sector can react to past policy decisions, giving incentives to the government to stick to its commitments. This is why the resulting reputation forces¹⁵ are often considered

¹⁵As in Ljungqvist and Sargent (2018), reputation forces refer to the incentives stemming from repeated interactions between a government and a continuum of atomistic private agents. In such macroeconomic games, individuals are not strategic but atomistic. Still, they may form history-dependent expectations that allow for trigger "strategies." In contrast with "reputation effects" in game theory, we do not assume any form of asymmetric information.

in the literature as an endogenous substitute for commitment ability, but they are also known to produce multiple equilibria, including in macroeconomic settings (see the initial contribution by [Barro and Gordon \(1983b\)](#)).

To study how reputation forces change the cost of controllability, we consider the repeated version of our macroeconomic game. In this repeated setting, deviating from the reaction function leads to a welfare cost for the government that measures its commitment ability, as in the static setting.

Compared with the static setting, only a larger commitment ability ensures the implementation of a unique equilibrium outcome. But why the multiplicity of equilibria due to dynamic incentives is more difficult to handle than in static settings? The higher cost of controllability results from reputation forces that make government's future pay-offs dependent on its response: due to history-dependent private reactions, the private sector can shift to an inferior continuation equilibrium when the government sticks to its commitment, and in contrast, to a better continuation equilibrium when the government deviates. Such history-dependent reactions lead to reputation forces for the government not to keep its commitment and to deviate from its reaction function in the form of a fixed cost. To get a unique equilibrium, the commitment ability must be sufficiently large to overcome these reputation forces in addition to the incentives that we identified in the static setting.

5.2 Interpreting κ ?

To conclude this section, let us discuss what may be behind the cost κ .

Reputation loss A first interpretation of the cost κ is that this cost captures the reputation loss associated with a deviation from past announcements (see [Dovis and Kirpalani \(forthcoming\)](#), for such a recent model or reputation loss in macroeconomics). Suppose, indeed, that there are two types of policymakers: one that can perfectly commit to reaction functions and the other one that cannot and sets its policy under discretion. When the type of the policymaker cannot be observed, private agents form beliefs about its type. In this environment, playing something else than the reaction function $\bar{y}(x)$ leads the discretion-type policymaker to reveal its type and to lose any gain from being pooled with the commitment-type policymaker.

REMARK. Whereas reputation loss may provide micro-foundations for commitment ability, we showed above that the repetition of the static game does not—the dynamic incentives in the repeated game setting is also sometimes called reputation in macroeconomics.

Political institutions A second interpretation of the cost κ is the cost due to the decision process needed to deviate from a reaction function. Such a deviation may require to change a law, a regulation or reversing the independence of an independent agency, requiring a potentially costly political and bureaucratic process, for example, as a result from disagreement between political parties (e.g., [Piguillem and Riboni \(2021\)](#), in the context of fiscal policy). Such kind of costs may exist when decisions are made by committees, as it is the case in public institutions such as central banks (see, e.g., [Riboni \(2010\)](#)).

Judicial institutions A third interpretation of the cost is the fact that legislations may put constraints on the degree of discretion of policymakers. This is, for example, the case of the Administrative Procedure Act of 1946 in the United States, which requires that regulations by agencies should not be “arbitrary and capricious, an abuse of discretion.” Especially under the “Hard Look” doctrine, this requires agencies to sufficiently motivate their decisions to change their set course of actions.

Intrinsic preference The potential embarrassment of the policymaker—when deviating from past commitments—emphasized by Woodford (2012) may be related to an intrinsic preference by the policymaker to stick to commitments. Also, such a preference may be shared by private agents so that a deviation by the policymaker may result into a welfare loss that the policymaker may internalize.

Cognitive cost/bounded rationality Finally, the cost associated with a deviation may simply be the cost of acquiring and processing information to find the optimal deviation.

6. CONCLUSION

In this paper, we investigate the ability of a government to implement a unique equilibrium outcome when its commitment ability is limited. We find that, surprisingly, implementation does not require large commitment ability in a relatively large set of static games. On top of the quantification of the commitment ability required for implementation, our results also give insights about the design of commitments, and especially, the importance of designing commitments that are ex post credible in and out of equilibrium. Interestingly, in general, designing credible rules is relatively simple as it simply relies on slight deviation from the ex post optimal policy when this ex post optimal policy is consistent with private decisions—a competitive outcome—and to follow the ex post optimal policy otherwise. Finally, we discuss potential limits to these results.

REFERENCES

- Albanesi, Stefania, Varadarajan V. Chari, and Lawrence J. Christiano (2003), “Expectation traps and monetary policy.” *Review of Economic Studies*, 70, 715–741. [1123]
- Armenter, Roc (2008), “A general theory (and some evidence) of expectation traps in monetary policy.” *Journal of Money, Credit and Banking*, 40, 867–895. [1123]
- Atkeson, Andrew, Varadarajan V. Chari, and Patrick J. Kehoe (2010), “Sophisticated monetary policies.” *Quarterly Journal of Economics*, 125, 47–89. [1123, 1134]
- Barro, Robert J. and David B. Gordon (1983a), “A positive theory of monetary policy in a natural rate model.” *Journal of Political Economy*, 91, 589–610. [1122]
- Barro, Robert J. and David B. Gordon (1983b), “Rules, discretion and reputation in a model of monetary policy.” *Journal of Monetary Economics*, 12, 101–121. [1122, 1123, 1146]

Barthélemy, Jean and Eric Mengus (2022), “Time-consistent implementation in macroeconomic games.” CEPR Discussion Papers 17120, C.E.P.R. Discussion Papers March 2022. [1122, 1123, 1138, 1144, 1145]

Bassetto, Marco (2005), “Equilibrium and government commitment.” *Journal of Economic Theory*, 124, 79–105. [1120, 1123, 1125, 1126, 1133, 1134, 1142]

Bassetto, Marco and Christopher Phelan (2008), “Tax riots.” *Review of Economic Studies*, 75, 649–669. [1123]

Bilbiie, Florin O. (2011), “The time inconsistency of delegation-based time inconsistency solutions in monetary policy.” *Journal of Optimization Theory and Applications*, 150, 657–674. [1123]

Chari, Varadarajan V., Lawrence J. Christiano, and Martin Eichenbaum (1998), “Expectation traps and discretion.” *Journal of Economic Theory*, 81, 462–492. [1123]

Chari, Varadarajan V. and Patrick J. Kehoe (1990), “Sustainable plans.” *Journal of Political Economy*, 98, 783–802. [1123, 1139, 1142]

Chari, Varadarajan V. and Patrick J. Kehoe (2016), “Bailouts, time inconsistency, and optimal regulation: A macroeconomic view.” *American Economic Review*, 106, 2458–2493. [1122, 1128]

Clarida, Richard, Jordi Galí, and Mark Gertler (2000), “Monetary policy rules and macroeconomic stability: Evidence and some theory.” *Quarterly Journal of Economics*, 115, 147–180. [1123]

Cochrane, John H. (2011), “Determinacy and identification with Taylor rules.” *Journal of Political Economy*, 119, 565–615. [1123]

Debortoli, Davide and Ricardo Nunes (2010), “Fiscal policy under loose commitment.” *Journal of Economic Theory*, 145, 1005–1032. [1123]

Dovis, Alessandro and Rishabh Kirpalani “Rules without commitment: Reputation and incentives.” *Review of Economic Studies*. (forthcoming). [1122, 1146]

Ennis, Huberto M. and Todd Keister (2009), “Bank runs and institutions: The perils of intervention.” *American Economic Review*, 99, 1588–1607. [1122]

Farhi, Emmanuel, Christopher Sleet, Ivan Werning, and Sevin Yeltekin (2012), “Non-linear capital taxation without commitment.” *The Review of Economic Studies*, 79, 1469–1493. [1123]

Farhi, Emmanuel and Jean Tirole (2012), “Collective moral hazard, maturity mismatch, and systemic bailouts.” *American Economic Review*, 102, 60–93. [1121, 1122, 1138, 1139, 1140, 1141]

Fischer, Stanley (1980), “Dynamic inconsistency, cooperation and the benevolent dissembling government.” *Journal of Economic Dynamics and Control*, 2, 93–107. [1142]

Gourinchas, Pierre-Olivier, Philippe Martin, and Todd E. Messer (2020), “The economics of sovereign debt, bailouts and the eurozone crisis.” Working Paper 27403, National Bureau of Economic Research. [1119]

Halac, Marina and Pierre Yared (2022), “Fiscal rules and discretion under limited enforcement.” *Econometrica*, 9, 2093–2127. [1123]

Halac, Marina and Pierre Yared “Instrument-based vs. Target-based rules.” Review of Economic Studies. (forthcoming). [1123]

Hall, Robert E. and Ricardo Reis (2016), “Achieving price stability by manipulating the central bank’s payment on reserves.” NBER Working Papers 22761, National Bureau of Economic Research, Inc. [1123]

Jensen, Henrik (1997), “Credibility of optimal monetary delegation.” *American Economic Review*, 87, 911–920. [1123]

Keister, Todd (2016), “Bailouts and financial fragility.” *Review of Economic Studies*, 83, 704–736. [1122]

King, Robert G. and Alexander L. Wolman (2004), “Monetary discretion, pricing complementarity, and dynamic multiple equilibria.” *Quarterly Journal of Economics*, 119, 1513–1553. [1123]

Kydland, Finn E. and Edward C. Prescott (1977), “Rules rather than discretion: The inconsistency of optimal plans.” *Journal of Political Economy*, 85, 473–491. [1122, 1123]

Ljungqvist, Lars and Thomas J. Sargent (2018), “Recursive macroeconomic theory, fourth edition number 0262038668.” In *MIT Press Books*, The MIT Press. [1126, 1145]

Loisel, Olivier (2009), “Bubble-free policy feedback rules.” *Journal of Economic Theory*, 144, 1521–1559. [1123]

Philippon, Thomas and Olivier Wang “Let the worst one fail: A credible solution to the too-big-to-fail conundrum.” 2021. Mimeo. [1122]

Piguillem, Facundo and Alessandro Riboni (2021), “Fiscal rules as bargaining chips.” *Review of Economic Studies*, 88, 2439–2478. [1146]

Riboni, Alessandro (2010), “Committees as substitutes for commitment.” *International Economic Review*, 51, 213–236. [1146]

Rogoff, Kenneth (1985), “The optimal degree of commitment to an intermediate monetary target.” *Quarterly Journal of Economics*, 100, 1169–1189. [1123]

Sargent, Thomas J. and Neil Wallace (1975), ““Rational” expectations, the optimal monetary instrument, and the optimal money supply rule.” *Journal of Political Economy*, 83, 241–254. [1123]

Schelling, Thomas C. (1960), *The Strategy of Conflict*. MIT Press Books, Harvard University Press. [1120]

Schneider, Martin and Aaron Tornell (2004), “Balance sheet effects, bailout guarantees and financial crises.” *Review of Economic Studies*, 07, 883–913. [1122]

Stokey, Nancy L. (1991), “Credible public policy.” *Journal of Economic Dynamics and Control*, 15, 627–656. [1126]

Taylor, John B. (1999), *Monetary Policy Rules*. NBER Books, National Bureau of Economic Research, Inc. [1123]

Walsh, Carl E. (1995), “Optimal contracts for central bankers.” *American Economic Review*, 85, 150–167. [1123]

Woodford, Michael (2012), “Methods of policy accommodation at the interest-rate lower bound.” Technical Report, Columbia University. [1120, 1147]

Co-editor Pierre-Olivier Weill handled this manuscript.

Manuscript received 20 September, 2023; final version accepted 18 January, 2024; available online 1 February, 2024.