

A Model of Repeated Collective Decisions*

Antonin Macé¹ and Rafael Treibich²

¹Paris School of Economics, CNRS and École Normale Supérieure-PSL

²University of Southern Denmark

March 16, 2023

Abstract

The theory of repeated games offers a compelling rationale for cooperation in a variety of environments. Yet, its consequences for collective decision-making have been largely unexplored. In this paper, we propose a general model of repeated voting in committees and study equilibrium behavior under alternative majority rules. Our main characterization reveals a complex relationship between the majority threshold, the preference distribution, and the set of equilibrium payoffs. In contrast with the stage game, equilibria in the repeated game may involve a form of implicit logroll, individuals sometimes voting against their preference to achieve the efficient decision. In turn, this affects the optimal voting rule, which may differ significantly from the optimal rule under sincere voting. The model provides a rationale for the use of unanimity rule, while accounting for the prevalence of consensus in committees which use a lower majority threshold.

*We are very grateful to Vincent Anesi, Geir Asheim and Alessandra Casella for useful comments, and to participants of the following seminars: University of Strasbourg, Paris School of Economics, Online Social Choice and Welfare, COMSOC, and of the following conferences: *Aggregation Across Disciplines*, Social Choice and Welfare, ACM Conference on Economics and Computation (EC'22).

1 Introduction

Collective decisions in many committees display two key features: (i) decisions are taken repeatedly over time, and (ii) committee members typically differ in how much they care about each decision. Examples include international organizations, city councils, hiring committees, etc. For instance, the Council of the European Union, one of the three main decision bodies of the EU, issues more than 300 legislative acts every year on a wide range of topics. Countries' stakes on these proposed reforms vary significantly depending on their economy, citizens' preferences, historical background, cultural norms, current national legislation, etc. In this paper, we propose a model of collective decisions that captures these two defining characteristics and study how they affect equilibrium voting behavior and the comparison of majority voting rules.

In the classical setting where the committee takes a single binary decision, sincere voting is a dominant strategy under any qualified majority rule. By contrast, in a repeated setting, voting behavior may be conditioned on past realizations and voting choices. This history dependence allows for increased cooperation as individuals may now enter into implicit logrolling agreements, sometimes voting against their preference to further the group's interest. As a result, equilibrium voting behavior need not be sincere and the evaluation of different majority rules may differ significantly from the static benchmark.

Our paper sheds light on two important facts about collective decision-making that can be hard to reconcile with standard voting theory. The first observation is that, despite strong similarities in their organization and structure, committees vary widely in the procedures and voting rules used for making collective decisions: from simple majority, unanimity rule, and weighted majority rules to the use of veto power. The heterogeneity of voting procedures can sometimes be observed within the same organization, decisions of different nature being taken according to different voting rules. For instance, the Council of the EU uses simple majority for some decisions, but the unanimity rule for others. Why do otherwise similar committees use different voting rules? And why is the unanimity rule so prevalent, when it is usually considered to be inefficient and prone to gridlock? The second

observation is that decisions in many committees are often made by consensus, without taking a formal vote, even when the committee is supposed to use a (possibly non-unanimous) voting rule (Urfalino, 2014). The prevalence of consensus may seem puzzling, especially in large committees, where preferences over decisions are rarely consensual. We show that both observations can be rationalized if we acknowledge the repeated nature of collective decisions.

In our model an ex-ante symmetric committee makes repeated binary decisions about whether to accept or reject proposed reforms. At each stage utilities are drawn and observed publicly. Utilities reflect individuals' cardinal preferences for the reform relative to the status quo. Individuals then vote simultaneously either in favor or against the reform and a collective decision is taken according to a qualified majority rule. We study the equilibrium outcomes of this repeated game.

Our first result characterizes the set of equilibrium payoffs. In the repeated game, the pivotal individual may be incentivized to vote against her preference today if the net reward from complying with prescribed behavior exceeds the loss from implementing the decision she dislikes. The extent to which such non-sincere behavior can be sustained at equilibrium, i.e. the power of intertemporal incentives, depends both on the discount factor and the difference between the largest and smallest equilibrium payoff, as they determine the largest admissible reward from complying with prescribed behavior. In turn, the power of intertemporal incentives affects which decision rule can be implemented today and at which (long-term) cost, driving both the largest and smallest equilibrium payoffs. The power of intertemporal incentives thus obtains as the solution to a fixed point equation. At the optimal equilibrium, the committee follows a stationary cooperation norm, adopting the efficient decision unless the pivotal voter disagrees with a stake larger than the power of intertemporal incentives. As individuals become more patient, the optimal equilibrium payoff increases, eventually reaching the first-best.

Our second result refines the previous characterization by focusing on individual voting behavior. We show that the cooperation norm of the optimal equilibrium allows for a substantial level of consensual decisions, even if preferences are never consensual. When the discount factor is large enough, any accepted reform can receive unanimous approval.

To explore the implications of our results on the comparison of voting rules, we then focus on a more stylized setting where utilities write as a sum of common and private components. We derive explicit formulas for the set of equilibrium payoffs under any possible majority rule and discount factor. This allows us to perform two comparative statics exercises. First, we characterize the optimal voting rule, the one which maximizes the largest equilibrium payoff. While simple majority is optimal in the static benchmark, a super-majority can be optimal in the repeated game. In fact, even the unanimity rule, which performs worst in the static game, may turn out to be optimal in the repeated game. Second, we show that, contrary to what we would obtain in a static model, the level of consensus may be higher when committee members have more heterogeneous preferences or when the majority threshold is lower.

The rest of the paper is organized as follows. [Section 2](#) presents a simple example to illustrate the main intuition behind our model. [Section 3](#) reviews the related literature. [Section 4](#) introduces the model and lays out the main assumptions. [Section 5](#) characterizes the set of equilibrium payoffs and the optimal consensus probability. [Section 6](#) revisits our general results on a more stylized model to perform comparative statics. [Section 7](#) discusses our main assumptions. [Section 8](#) concludes.

2 An Example

A group of three individuals decides whether to accept or reject repeated proposals at either *unanimity* or *simple majority*. Assume proposals are of the following two types:

- Proposals of type A benefit one individual strongly ($u_i = 5$), one individual weakly ($u_i = 1$), and hurt one individual mildly ($u_i = -3$).
- Proposals of type B benefit one individual strongly ($u_i = 5$) and hurt two individuals weakly ($u_i = -1$).

Individuals are equally likely to occupy any position for each kind of proposals. We denote by $p \in (0, 1)$ the probability of occurrence of proposals A. Since both

proposals generate an average utility of 1, the first-best consists in accepting any proposal and yields ex-ante utility 1.

If the organization only makes one decision, voting sincerely is a weakly dominant strategy. Under *unanimity*, no reform is ever accepted at equilibrium, yielding ex-ante utility 0. Under *simple majority*, only proposals A are accepted at equilibrium, yielding ex-ante utility $p > 0$. Simple majority dominates unanimity but does not achieve the first best.

If the organization makes (infinitely) repeated decisions, it is possible to incentivize individuals to sometimes vote against their preference so as to achieve a higher level of utility. Consider the following strategy profile under *unanimity*: on the equilibrium path all three individuals always vote in favor of proposals A and B; any deviation is punished by a permanent reversal to the stage-game equilibrium. Such a profile is an equilibrium if and only if the worst-off individual in proposals A has an incentive to vote in favor of the proposal. This is the case if her inter-temporal utility from the reform being adopted, $-3(1 - \delta) + \delta$, exceeds her inter-temporal utility from the reform being rejected, 0. As a result, the first-best can be sustained under this profile if $\delta \geq 3/4$. In fact, this equilibrium is optimal, and the previous condition is thus necessary for the first-best to be achieved at equilibrium.

Under *simple majority*, sustaining the first-best no longer requires incentivizing A's worst-off individual to vote in favor since proposals A are accepted under sincere voting. Consider the following strategy profile: on the equilibrium path all three individuals vote in favor of proposals B, but they vote sincerely on proposals A; any deviation is punished by a permanent reversal to the stage-game equilibrium. This profile is an equilibrium if and only if $-(1 - \delta) + \delta \geq \delta p$, or equivalently $\delta \geq 1/(2 - p)$. This equilibrium is optimal, and the previous condition is thus necessary for the first-best to be achieved at equilibrium.¹ In contrast to unanimity, achieving efficiency under simple majority only requires incentivizing individuals with a stake of 1, as opposed to 3, and may thus seem easier to achieve. However,

¹Under simple majority, the payoff of the stage-game equilibrium is not the lowest feasible payoff ($p > 0$). However, it can be shown (in this example) that the worst equilibrium payoff in the repeated game does in fact coincide with the payoff of the stage-game equilibrium.

since simple majority outperforms unanimity in the stage game, the long term benefit from complying with prescribed behavior is smaller, which can eventually make efficiency harder to sustain. Here, if $p > 2/3$, there exists a range of discount factors, $\delta \in (3/4, 1/(2 - p))$, where only unanimity can sustain the first-best.

3 Literature Review

Most of the literature on repeated collective decision-making considers decisions over a single, persistent, issue (e.g. a central bank setting the interest rate). The objective is often to understand how the endogeneity of the status quo (today’s decision becomes tomorrow’s status quo) affects voting behavior, most often under the assumption of Markovian strategies. For instance, [Baron \(1996\)](#) proposes a model of dynamic redistributive politics where legislators have single peaked preferences over a one-dimensional policy space. He characterizes voting behavior at the stationary equilibrium and show how it relates to the preference of the median voter. More recently, [Dzuida and Loeper \(2018\)](#) study an infinite horizon model where a group of legislators repeatedly make the same binary decision. They show how the endogeneity of the status quo leads to more polarized voting behavior at equilibrium than under sincere voting.² By contrast, our model considers a committee making repeated collective decisions about exogenous and independent proposals, reflecting committees with multiple functions and responsibilities. The repeated, as opposed to dynamic, nature of our model allows us to consider history dependent strategies, which are essential to generate the implicit logrolling behavior characteristic of our extremal equilibria. While the benefit of repetition on cooperation has long been acknowledged in the literature on repeated games ([Mailath and Samuelson, 2006](#)) and relational contracts more broadly ([Levin, 2003](#)), very few papers have explored the implications for voting behavior and optimal voting

²Other references in this literature include [Kalandrakis \(2004\)](#), who extends [Baron \(1996\)](#)’s model of legislative bargaining to a dynamic setting by considering a three players version of the divide the dollar game with endogenous reversion point, [Duggan and Kalandrakis \(2012\)](#), who prove the existence of stationary Markov perfect equilibria in an infinite-horizon model of legislative policy making with endogenous status quo, under general assumptions about the policy space and preferences, and [Penn \(2009\)](#), who studies the formation of “far-sighted” preferences under endogenous status quo when voters must choose, at each stage, between a randomly drawn proposal and the status quo.

rules. This is the main purpose of our paper.

Going back to [Buchanan and Tullock \(1962\)](#), it has been argued that making multiple, as opposed to a single, collective decisions opens up the possibility of vote trading between committee members. As noted in [Buchanan and Tullock \(1962\)](#), “The existence of a logrolling process is central to our general analysis of simple majority voting”. The literature on logrolling usually considers the possibility of vote trading over a finite and explicit set of alternatives under complete information ([Park, 1967](#); [Casella and Palfrey, 2019](#); [Casella and Macé, 2021](#)). Our model does not assume explicit vote trading agreements between voters. Instead, logrolling emerges endogenously at equilibrium, as an implicit agreement between voters. Increased cooperation can be sustained at equilibrium because of the threat of reverting to an inefficient equilibrium in case of failure to achieve the efficient decision.

A few papers from the logrolling literature are particularly connected to our analysis. In the experiments reported in [Fischbacher and Schudy \(2014, 2020\)](#), voters that have benefited from non-sincere voting from other voters in the past reciprocate when given the opportunity. In turn, trusting voters may vote non-sincerely at the beginning of the sequence, anticipating that their fellow committee members will reciprocate. This feature is also present in [Carrubba and Volden \(2000\)](#), who consider a sequential model of logrolling embedded in a repeated model of coalition formation in a legislature. At each period a winning coalition of legislators is picked endogenously and a series of proposals, each benefiting only one of the coalition’s members (at the expense of all other legislators), is put to a vote. The objective of the model is to show why legislators may form oversized coalitions (larger than the majority threshold) to enter into logrolling agreements. The result relies on the legislators having a smaller incentive to deviate once their own bill has been passed if rejecting a proposal requires the defection of more than one legislator in the coalition. [Charroin and Vanberg \(2021\)](#) compare simple majority and unanimity rule in a model of logrolling with three voters. Relying on simulations and lab experiments, they obtain that logrolling is welfare-improving under unanimity but not necessarily under simple majority and that unanimity may outperform simple majority. Our results show that such comparative statics may hold

analytically for any number of voters and any super-majority requirement.

Finally, the paper closest to ours is [Maggi and Morelli \(2006\)](#), who study self-enforcing voting rules in international organizations. A group of countries takes repeated collective actions under the assumption that an action is only effective if taken by all countries. At the optimal equilibrium, countries take action (unanimously) if the number of countries in favor exceeds a fixed threshold. The decision rule on the equilibrium path thus mimics the outcome of a (super)majority rule under enforceable decisions. By contrast, the voting rule in our model matters in an explicit way, and a winning coalition can impose its preference on the losing minority. Similar to [Maggi and Morelli \(2006\)](#), we find that it is possible to achieve the first-best if the discount factor is large enough. However, our analysis of the equilibrium payoffs differs substantially. In particular, we consider a model where collective action does not require unanimous implementation and members' preferences have varying intensities. This latter aspect is crucial in many environments and greatly exacerbates the benefits of logrolling.

4 Setup

4.1 Stage game

Model - A group of individuals $N = \{1, \dots, n\}$ must collectively decide whether to adopt a proposed *reform* or keep the status quo. If enacted, the reform yields utility $u_i \in \mathbb{R}$ to individual $i \in N$, while the status quo's utility is normalized to 0. An individual $i \in N$ is thus *favorable* to the reform if $u_i \geq 0$, *opposed* if $u_i < 0$, and we refer to $|u_i|$ as her *stake* in the collective decision. The reform is characterized by the utility vector $\mathbf{u} = (u_i)_{i \in N}$, with mean $\bar{u} = (\sum_{i \in N} u_i)/n$. We say that a reform is *good* when $\bar{u} \geq 0$ and *bad* when $\bar{u} < 0$.

The reform $\mathbf{u}$ is drawn from a cumulative distribution function F , whose support is denoted by $S \subseteq \mathbb{R}^N$. We make the following three assumptions on F .

Assumption 1. F is symmetric: for any $\mathbf{u} \in S$ and any permutation π of N , $F(\mathbf{u}) = F(\mathbf{u}_\pi)$.

Assumption 2. F is smooth and its support S is bounded and convex.

Assumption 3. $\mathbb{E}_F[|\bar{u}|] < +\infty$.

Assumption 1 reflects the fact that individuals are ex-ante identical, but ex-post heterogeneous. **Assumption 2** is made for ease of exposition, it ensures in particular that the distribution F has no atom. Finally, **Assumption 3** guarantees that the game has well-defined expected payoffs.

After observing reform $\mathbf{u}$, each individual i votes either in favor or against the reform. A *majority rule* with threshold $m \in [1/2, 1]$ then decides the reform's fate $d \in \{0, 1\}$. If at least mn individuals vote in favor, the reform is adopted, $d = 1$, and each individual i gets utility u_i . If not, the status quo remains, $d = 0$, and each individual gets utility 0.

Strategies and Equilibrium - A *voting strategy* v_i for player i associates to any reform $\mathbf{u}$ a vote $v_i(\mathbf{u}) \in \{0, 1\}$. At the unique Nash equilibrium in weakly undominated strategies,³ every individual votes sincerely, i.e. $v_i(\mathbf{u}) = \mathbb{1}\{u_i \geq 0\}$. As a result, a reform is collectively accepted if and only if the individual with the $[mn]$ -th largest utility is favorable, i.e. $u_{[mn]} \geq 0$. This individual plays a central role in our analysis as she is a *pivot*, both in the static and in the repeated game, and she is thus denoted by p . Henceforth, we refer to her utility $u_p := u_{[mn]}$ as the *pivot utility*. We note that the equilibrium is usually inefficient, as bad reforms may be accepted (when $u_p \geq 0$ but $\bar{u} < 0$) and good reform may be rejected (when $u_p \leq 0$ but $\bar{u} > 0$) under sincere voting.

Decision rules - Given majority rule m , a strategy profile $(v_i)_{i \in N}$ induces a (group) decision rule. Formally, a *decision rule* d associates to any reform $\mathbf{u}$ a collective decision $d(\mathbf{u}) \in \{0, 1\}$. Throughout, we focus on symmetric decision rules, and we denote the (common) expected utility attached to such a rule d by⁴

$$U(d) = \int u_i d(\mathbf{u}) dF(\mathbf{u}). \quad (1)$$

³Individuals with utility $u_i = 0$ are indifferent between voting in favor or against, but this instance almost never arises under **Assumption 2**. The equilibrium in undominated strategies is thus *essentially* unique.

⁴The formula does not depend on the identity of $i \in N$ since both d and F are assumed to be symmetric.

We denote by d^0 the *sincere* decision rule induced by the stage-game equilibrium, $d^0(\mathbf{u}) = \mathbb{1}\{u_p \geq 0\}$.

Cost of Implementation - For any reform $\mathbf{u} \in S$ and any collective decision $d \in \{0, 1\}$, let $c(\mathbf{u}, d)$ denote the smallest transfer that would make the pivot favor decision d ,

$$c(\mathbf{u}, d) = \begin{cases} 0 & \text{if } d = d^0(\mathbf{u}) \\ |u_p| & \text{if } d \neq d^0(\mathbf{u}). \end{cases}$$

We refer to $c(\mathbf{u}, d)$ as the *cost of implementing* decision d on reform $\mathbf{u}$. It is equal to 0 if the pivot favors d and to the *pivot stake* $|u_p|$ if she does not. The cost $c(\mathbf{u}, d)$ reflects the smallest common reward needed to incentivize the required majority of m individuals to favor decision d on reform $\mathbf{u}$.⁵ For any decision rule $d(\cdot)$, we define $C(d)$ and $\Delta(d)$ as, respectively, the expected and the largest cost of implementing decision $d(\mathbf{u})$,

$$C(d) = \mathbb{E}[c(\mathbf{u}, d(\mathbf{u}))] \quad \text{and} \quad \Delta(d) = \sup_{\mathbf{u} \in S} c(\mathbf{u}, d(\mathbf{u})).$$

Note that $C(d^0) = \Delta(d^0) = 0$ since the pivot always agrees with d^0 .

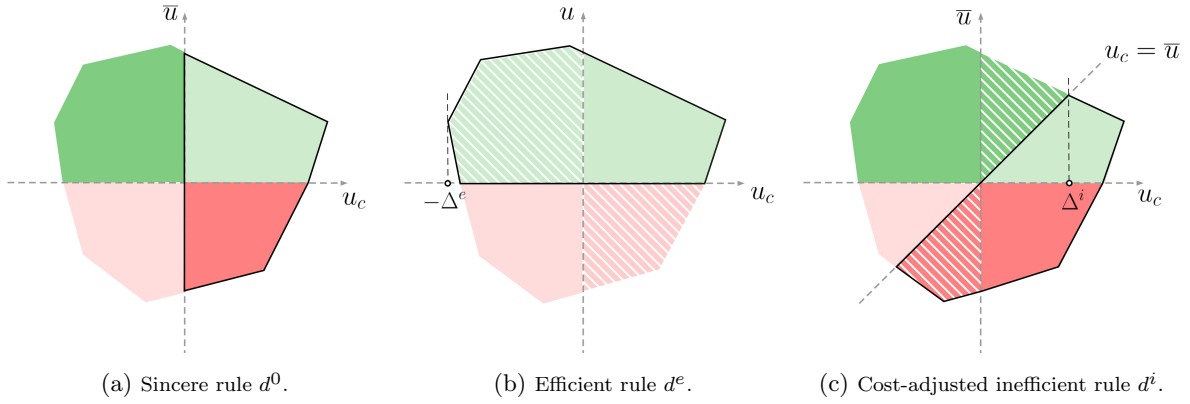
Two benchmark rules - Two decision rules play an important role in our characterization of the repeated game's equilibrium payoffs. First, we define the *efficient* rule d^e as the rule which accepts good reforms and rejects bad reforms, i.e. $d^e(\mathbf{u}) = \mathbb{1}\{\bar{u} \geq 0\}$. In what follows, we focus on the interesting case where the stage-game is inefficient, i.e. $U^e := U(d^e) > U^0 := U(d^0)$.

Second, we define the *cost-adjusted inefficient rule* d^i as the rule which selects the inefficient decision if and only if the inefficiency exceeds the cost of implementing the corresponding decision: $d^i(\mathbf{u}) = 1 - d^e(\mathbf{u})$ if and only if $c(\mathbf{u}, 1 - d^e(\mathbf{u})) \leq |\bar{u}|$. We note that d^e maximizes $U(\cdot)$ over all symmetric rules, while d^i minimizes $U(\cdot) + C(\cdot)$ over all symmetric rules. In the sequel, we denote the largest costs of implementation of the two benchmark rules by $\Delta^e = \Delta(d^e)$ and $\Delta^i = \Delta(d^i)$.

⁵The focus on a common (rather than personalized) reward is a consequence of [Assumption 5](#) below.

We illustrate decision rules d^0 , d^e and d^i on Figure 1. A reform is identified graphically by the average utility $\bar{u}$ and the pivot utility u_p . The colored area represents the (projected) support S of F . Good reforms are represented in green, while bad reforms are represented in red. Collectively accepted proposals are located inside the black polygon. Inefficient decisions, either proposals inefficiently rejected or inefficiently accepted, are represented in darker tones. Reforms in hatched areas have a positive cost of implementation, i.e. $c(\mathbf{u}, d(\mathbf{u})) = |u_p| > 0$, while reforms outside of hatched areas have a zero cost of implementation.

Figure 1: Three decision rules.



4.2 Repeated Game

We now consider an infinitely repeated version of the stage game.

Timing - At each stage, a reform $\mathbf{u}$ is drawn from F independently of previous stages. The reform $\mathbf{u}$ is publicly observed, then individuals simultaneously vote under majority rule m to decide the reform's fate. The *history* at time t , denoted by h^t , consists of all reforms and votes prior to t . A *strategy* σ_i for individual i associates to every history h and reform $\mathbf{u} \in S$ a vote $\sigma_i(h, \mathbf{u}) \in \{0, 1\}$. Utilities are discounted with a discount factor $\delta \in (0, 1)$. We say that a history h is *on the path* of a strategy profile $\sigma = (\sigma_i)_{i \in N}$ if the votes at each period are the ones

specified by σ given the utility realizations.⁶

Assumptions - We restrict our analysis to strategy profiles that satisfy the following three properties.

Assumption 4 (Symmetry). *All individuals use the same strategy.*

Assumption 5 (Independence of Individual Votes (IIV)). *Voting behavior only depends on past anonymized realizations of utilities and decisions.*

Assumption 6 (As-if-pivotal Voting). *Taking as given future play prescribed by the strategy profile, players play as if they were pivotal.*

Symmetry is a natural assumption in our context since the model is fully symmetric ex-ante. IIV requires that strategies only depend on past decisions, and not on past individual votes. This implies that deviations from the equilibrium path may only be punished collectively. IIV can be justified by a reluctance to single out or antagonize specific individuals for their votes.⁷ Similar requirements have been used in the literature to allow for history-dependent strategies, without losing too much tractability (Bernheim and Slavov, 2009; Anesi and Seidmann, 2015). As-if-pivotal Voting allows us to pin down the voting behavior of individuals who are not pivotal, which is left unconstrained by subgame perfection. Individuals then vote for the alternative that maximizes their continuation utility. The requirement is commonly used in dynamic voting games (see for instance Ali et al. (2022)), it was first proposed in Baron and Kalai (1993) as *stage-undomination*.⁸ We discuss in more detail these assumptions in Section 7.

Equilibrium - We focus on subgame perfect equilibria that satisfy the above three assumptions, and simply refer to them as *equilibria*. Denoting by $\underline{w}_\delta$ and $\bar{w}_\delta$ the

⁶Such histories are also referred to in the literature as consistent histories (Mailath and Samuelson (2006)).

⁷In fact, all of our results hold if we weaken IIV to the requirement of *Anonymity*, i.e. that voting behavior only depends on past anonymized realizations of utilities and votes (voting behavior may then be conditioned on the fraction of individuals voting in favor of the reform). We require the slightly stronger property for ease of exposition.

⁸Note that when strategies do not depend on the detail of past individual votes, but only on the resulting collective decision (as under IIV), imposing As-if-pivotal Voting does not affect the set of equilibrium payoffs.

lowest and highest equilibrium payoffs,⁹ respectively, we define,

$$\kappa_\delta := \frac{\delta}{1-\delta} (\bar{w}_\delta - \underline{w}_\delta) \quad (2)$$

as the *power of intertemporal incentives*. Parameter κ_δ reflects how costly a decision can be implemented at equilibrium in the first stage, given that continuation promises must themselves be equilibrium payoffs and future is discounted by δ .

5 Equilibrium

We start by introducing truncated decision rules, which play an important role in our characterization of equilibrium payoffs.

5.1 Truncated Rules

We have highlighted in the previous section that decision rules sometimes involve positive implementation costs, which might be prohibitive in the repeated game if intertemporal incentives are not powerful enough. This consideration leads us to the following definition. For any decision rule d and any cutoff $\kappa \in \mathbb{R}_+$, the *truncated rule* d_κ selects decision $d(\mathbf{u})$ unless its cost of implementation is greater than κ , in which case it selects the sincere decision,

$$d_\kappa(\mathbf{u}) = \begin{cases} d(\mathbf{u}) & \text{if } c(\mathbf{u}, d(\mathbf{u})) \leq \kappa \\ d^0(\mathbf{u}) & \text{if } c(\mathbf{u}, d(\mathbf{u})) > \kappa. \end{cases}$$

For any cutoff $\kappa \in \mathbb{R}_+$, let $\mathcal{D}_\kappa$ denote the set of all decision rules truncated by κ . A decision rule can be implemented at equilibrium in the first stage if and only if it belongs to $\mathcal{D}_{\kappa_\delta}$. Indeed, as long as the cost of implementation is smaller than κ_δ , the pivot may be incentivized to vote in any desired direction by rewarding compliance with continuation promise $\bar{w}_\delta$ and punishing deviations with continuation promise $\underline{w}_\delta$.

Two truncated rules play a central role in our characterization of equilibrium

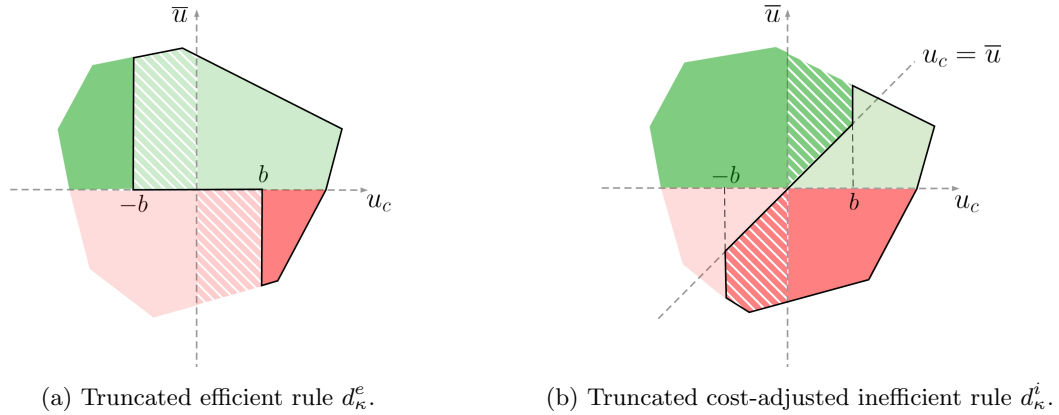
⁹The set of equilibrium payoffs is a compact interval, see Lemmas 1 to 3.

payoffs. Under the *truncated efficient rule* d_κ^e , the decision taken is efficient unless its cost of implementation exceeds κ . That is, good reforms are accepted unless the pivot utility is too small ($u_p \leq -\kappa$), while bad reforms are rejected unless the pivot utility is too high ($u_p \geq \kappa$). When $\kappa = 0$, the rule d_κ^e coincides with the sincere rule d^0 as no decision with positive implementation cost is allowed. For $\kappa \in (0, \Delta^e)$, the rule d_κ^e coincides with the sincere rule whenever it prescribes efficient decisions, while reversing some, but not all, inefficient decisions. For $\kappa \geq \Delta^e$, the truncation becomes moot (since Δ^e is the highest implementation cost of an efficient decision) and d_κ^e coincides with the efficient rule. In terms of utility, $U(d_\kappa^e)$ is increasing in κ , going from U^0 when $\kappa = 0$ to the first best U^e when $\kappa \geq \Delta^e$.

Under the *truncated cost-adjusted inefficient rule* d_κ^i , the decision taken on reform $\mathbf{u}$ is inefficient unless its cost of implementation exceeds $\bar{u}$ or exceeds κ . That is, good reforms are rejected unless the pivot utility is too high ($u_p \geq \min(\bar{u}, \kappa)$), while bad reforms are accepted unless the pivot utility is too low ($u_p \leq \max(\bar{u}, -\kappa)$). As the cutoff κ increases, more inefficient decisions are taken. The utility $U(d_\kappa^i)$ thus decreases with κ , and the rule d_κ^i coincides with d^i whenever $\kappa \geq \Delta^i$.

We illustrate the two truncated rules on [Figure 2](#) below.

Figure 2: Truncated decision rules.



5.2 Characterization of equilibrium payoffs

For any cutoff $\kappa \in \mathbb{R}_+$, let $\bar{V}(\kappa)$ and $\underline{V}(\kappa)$ be such that,

$$\bar{V}(\kappa) := \max_{d \in \mathcal{D}_\kappa} U(d) \quad \text{and} \quad \underline{V}(\kappa) := \min_{d \in \mathcal{D}_\kappa} (U(d) + C(d)).$$

In words, $\bar{V}(\kappa)$ is the highest payoff achievable by a κ -truncated rule, while $\underline{V}(\kappa)$ is the lowest cost-augmented payoff achievable by a κ -truncated rule. Since both optimizations can be performed reform by reform, it follows that $\bar{V}(\kappa)$ and $\underline{V}(\kappa)$ can be implemented, respectively, by the truncated efficient rule d_κ^e and the truncated cost-adjusted inefficient rule d_κ^i ,

$$\bar{V}(\kappa) = U(d_\kappa^e) \quad \text{and} \quad \underline{V}(\kappa) = U(d_\kappa^i) + C(d_\kappa^i).$$

Furthermore, because set $\mathcal{D}_\kappa$ expands when κ increases, we get that $\bar{V}(\kappa)$ increases with κ , while $\underline{V}(\kappa)$ decreases with κ .

Theorem 1. *The set of equilibrium payoff is equal to $[\underline{V}(\kappa_\delta), \bar{V}(\kappa_\delta)]$, where the power of intertemporal incentives κ_δ is given by:*

$$\kappa_\delta = \max \left\{ \kappa \geq 0 \mid (1 - \delta)\kappa = \delta [\bar{V}(\kappa) - \underline{V}(\kappa)] \right\}. \quad (3)$$

Theorem 1 characterizes the set of equilibrium payoffs in the repeated game through the power of intertemporal incentives κ_δ , defined as a solution to fixed-point equation:

$$(1 - \delta)\kappa_\delta = \delta [\bar{V}(\kappa_\delta) - \underline{V}(\kappa_\delta)]. \quad (4)$$

The intuition for **Theorem 1** follows from the dual relationship between the power of intertemporal incentives κ_δ and the extremal equilibrium payoffs $\underline{w}_\delta$ and $\bar{w}_\delta$.¹⁰ First, the power κ_δ can be simply expressed from the extremal payoffs, as in (2). Second, to see how extremal payoffs depend on power κ_δ , recall that any decision rule in $\mathcal{D}_{\kappa_\delta}$ can be implemented at the first stage of the repeated game. Now, the key observation is that there is no on-path cost at incentivizing high payoffs,

¹⁰The proof is provided in the appendix, it relies on the technique of self-generation, see for instance [Mailath and Samuelson \(2006\)](#).

since high payoffs can be sustained by high rewards (on-path) and low punishments payoffs (off-path). Hence, the highest equilibrium payoff that can then be sustained is simply given by the payoff-maximizing decision rule in $\mathcal{D}_{\kappa_\delta}$, that is $\bar{w}_\delta = \bar{V}(\kappa_\delta) = U(d_{\kappa_\delta}^e)$.

By contrast, there may be on-path costs at incentivizing low payoffs, since low payoffs may require high rewards which are then realized on-path. As a result of this trade-off, the lowest equilibrium payoff is obtained by minimizing the cost-augmented payoff $U(\cdot) + C(\cdot)$ on $\mathcal{D}_{\kappa_\delta}$, and we obtain $\underline{w}_\delta = \underline{V}(\kappa_\delta) = U(d_{\kappa_\delta}^i) + C(d_{\kappa_\delta}^i)$. This intertemporal utility can be achieved by a non-stationary equilibrium which implements $d_{\kappa_\delta}^i$ at the first stage. Intuitively, to achieve the lowest equilibrium payoff, an inefficient decision is taken at the first stage only if its inefficiency exceeds its implementation cost, which then realizes on the equilibrium path.

5.3 Equilibrium outcomes

In this section, we derive a number of implications of [Theorem 1](#) on equilibrium outcomes.

Corollary 1. *The optimal equilibrium payoff can be achieved by a stationary equilibrium in which the truncated efficient rule $d_{\kappa_\delta}^e$ is implemented at each stage.*

[Corollary 1](#) implies that the optimal equilibrium may be understood as a (self-enforcing) *cooperation norm*,¹¹ whereby voters routinely abide by collectively efficient decisions unless the pivot stake exceeds cutoff κ_δ . In turn, the parameter κ_δ may be interpreted as the optimal *degree of cooperation* that is achievable in the repeated game.

Corollary 2. *The optimal degree of cooperation κ_δ is weakly increasing with δ .*

A consequence of [Corollary 2](#) is that the set of equilibrium payoffs expands with δ . Then, with a higher discount factor δ , a higher degree of cooperation can be achieved for two reasons: because voters attach a greater importance to the future (relative to the present) and because the available rewards and punishments become richer.

¹¹In fact, this stationary path constitutes the unique optimal equilibrium path.

The next result addresses the following questions: Can even the first best be achieved? If not, can there be any cooperation? We say that *partial cooperation* is achieved when an equilibrium delivers a utility strictly higher than U^0 . We say that *full cooperation* is achieved when an equilibrium delivers the first-best utility U^e .

Corollary 3. *There are thresholds δ^P, δ^F with $0 < \delta^P \leq \delta^F < 1$ such that:*

- *if $\delta < \delta^P$, the unique equilibrium payoff is U^0 , the stage-game equilibrium payoff,*
- *if $\delta \in [\delta^P, \delta^F)$, only partial cooperation is possible at equilibrium,*
- *if $\delta \geq \delta^F$, full cooperation is possible at equilibrium.*

Corollary 3 characterizes when either partial or full cooperation can be achieved at equilibrium. For intermediate values of the discount factor, the optimal cooperation norm coincides with a (strictly) truncated efficient rule. For higher values of δ , even the efficient rule can be sustained.

5.4 Consensus

While the previous section focuses on equilibrium outcomes at the group level, we now derive implications of **Theorem 1** for individual voting behavior. We focus on optimal equilibria and ask how much *consensus* they can generate.

For a stationary equilibrium σ , we let v^σ denote the voting profile used at each stage. The *consensus probability* under σ is then defined as the probability that a reform is approved unanimously, that is $\mathbb{P}(\prod_{i \in N} v_i^\sigma(\mathbf{u}) = 1)$. We focus on the *optimal consensus probability* which is defined as the highest consensus probability that can be achieved at an optimal equilibrium: if Σ_δ is the set of stationary optimal equilibria, we define:¹²

$$P_\delta = \max_{\sigma \in \Sigma_\delta} \mathbb{P} \left(\prod_{i \in N} v_i^\sigma(\mathbf{u}) = 1 \right).$$

¹²We note that the highest consensus probability at an equilibrium of the repeated game might be attained for non-optimal equilibria. Here, we treat consensus as a secondary objective that comes after payoff maximization. While not exhaustive, we believe that this approach still provides a reasonable and tractable benchmark to highlight how consensus can emerge in the repeated game.

Proposition 1. *The optimal consensus probability P_δ is weakly increasing with δ . There is $\delta^C \in (0, 1)$ such that the optimal consensus probability is maximal for $\delta \geq \delta^C$. All good reforms are then accepted with consensus.*

When the discount factor is higher, the higher power of intertemporal incentives allows to construct optimal equilibria with worse punishments. The increase in the optimal consensus probability thus reflects two effects. First, the optimal degree of cooperation is higher and thus more good reforms are approved.¹³ Second, any good reform is more likely to receive consensus, as harsher punishments can incentivize even the least favorable voter to vote in favor of the reform.

The first substantive lesson of [Proposition 1](#) is that consensus becomes easier to achieve when voters become more patient. This refines the observation made in [Corollary 2](#) on self-enforcing cooperation norms. When δ increases, it becomes possible to implement cooperation norms such that collective decisions are better on average and also such that reforms are more likely to be approved unanimously. Moreover, if the discount factor is high enough, i.e. $\delta \geq \max(\delta^F, \delta^C)$, a cooperation norm can be sustained whereby only good reforms are approved, and all are approved unanimously.

6 Comparative statics

In this section, we derive rich comparative statics in a simplified model. We focus on a stylized class of preference distributions where (i) individual utilities write as the weighted sum of a common component and a private bias, and (ii) the distribution of biases is fixed across reforms.

6.1 Preference distributions

We assume that the utility of each individual $i \in N$ now writes as,

$$u_i = \theta + \alpha\beta_i$$

¹³This also means that less bad reforms are approved, but as we note in the proof, these reforms never receive a consensual approval.

where (i) the ordered vector of biases is fixed across reforms: there exist $b_1 \leq \dots \leq b_n$ such that $\beta_{[j]} = b_j$ for all $j \in \{1, \dots, n\}$, and (ii) the common component θ is drawn uniformly on $[-1, 1]$. As in the general model, we assume that the distribution of $\mathbf{u} = (u_i)_{i \in N}$ is symmetric, so that any voter has an equal chance to have a bias at any given rank of the distribution $\mathbf{b} = (b_i)_{i \in N}$. Parameter α reflects the (relative) weight attached to the bias, which we interpret as a measure of *diversity* within the committee. For ease of exposition, we assume that $b_1 = -1$, that $b_{n/2} < \bar{b} = 0$ and $\alpha \in [0, 1]$.¹⁴

In this setting, the difference between the average and the pivot utility is constant, as $\bar{u} - u_p = -\alpha b_p$, while the average utility, $\bar{u} = \theta$, varies across reforms. This feature greatly simplifies the computation of the optimal equilibrium, allowing us to study the effects of the majority threshold and of the degree of diversity on equilibrium outcomes.

A proposal is good if and only if the common component is positive, $\theta \geq 0$, while it is collectively accepted at the stage-game equilibrium if and only if the pivot is in favor, $\theta \geq -\alpha b_p$. As $b_{n/2} < 0$, we have $b_p \leq b_{n/2} < 0$ for any majority rule, so that the inefficiency of the stage-game equilibrium comes from good reforms $\theta \in [0, -\alpha b_p]$ being rejected. Under any voting rule m , the expected payoff under the efficient and sincere decision rules are now equal to:

$$U^e = \frac{1}{4} \quad \text{and} \quad U_m^0 = U^e - \frac{(\Delta_m^e)^2}{4}, \quad (5)$$

where $\Delta_m^e = \alpha |b_p| = \alpha |b_{(1-m)n}|$. Here, the parameter Δ_m^e can be interpreted as a measure of the *inefficiency* associated with majority rule m in the stage game. In particular, we get that the equilibrium utility in the stage-game is decreasing with threshold m .

Proposition 2. *In the stage game of the simplified model, the best majority rule is simple majority ($m = 1/2$) and the worst majority rule is unanimity ($m = 1$).*

We note that, although majority rule is optimal, it does not achieve the first best since $\Delta_{1/2}^e = |\alpha b_{n/2}| > 0$ and thus $U_{1/2}^0 < U^e$.

¹⁴When $\alpha \in [0, 1]$, there are always some reform that are good (resp. bad) for all voters, this is not essential for the results but it slightly simplifies the analysis.

6.2 Optimal voting rule

We turn to the analysis of the repeated game. As we show in the appendix, by applying [Theorem 1](#) for each possible value of the inefficiency Δ^e , we can obtain closed-form formulas for the optimal degree of cooperation κ_δ . In turn, we obtain the partial and full cooperation thresholds δ^F and δ^P , as well as the optimal utility $\bar{w}_\delta$. This eventually allows us to characterize the optimal voting rule, i.e. the rule which maximizes the highest equilibrium payoff $\bar{w}_\delta$.

Theorem 2. *In the simplified model, there are thresholds $\underline{\delta}, \bar{\delta}$ with $0 < \underline{\delta} < \bar{\delta} < 1$ and continuous functions $m^*, m^{**} : [0, 1] \rightarrow [1/2, 1]$ such that:*

- *if $\delta < \underline{\delta}$, the optimal rule is simple majority ($m = 1/2$)*
- *if $\underline{\delta} \leq \delta < \bar{\delta}$, the optimal rule is a (super-) majority $m^*(\delta)$*
- *if $\delta \geq \bar{\delta}$, any majority rule m with $m^*(\delta) \leq m \leq m^{**}(\delta)$ is optimal.*

Moreover, $m^*(\cdot)$ is weakly decreasing, from $m^*(\underline{\delta}) > 1/2$ up to $m^*(1) = 1/2$, while $m^{**}(\cdot)$ is weakly increasing, from $m^{**}(\bar{\delta}) = m^*(\bar{\delta})$ up to $m^{**}(1) = 1$.

[Theorem 2](#) shows that accounting for the repetition of collective decisions can drastically affect the assessment of alternative majority rules. In the benchmark case where the discount factor is low ($\delta < \underline{\delta}$), future is so discounted that the stage-game result of [Proposition 2](#) applies, i.e. simple majority ($m = 1/2$) is the optimal rule. In the opposite benchmark where the discount factor is high ($\delta \geq \bar{\delta}$), intertemporal incentives are sufficient to achieve full cooperation for some rules ($m^*(\delta) \leq m \leq m^{**}(\delta)$), and eventually for all rules when δ is large enough. This last remark coincides with the standard folk theorem.

The most interesting case of [Theorem 2](#) arises for intermediate discount factors ($\underline{\delta} \leq \delta < \bar{\delta}$). The optimal rule is then a majority m^* that is a strict supermajority at first ($m^*(\underline{\delta}) > 1/2$) and then weakly decreases for higher discount factors. In particular, this shows that a strict supermajority can dominate simple majority when accounting for the repetition of collective decisions, while the opposite conclusion would be drawn by focusing only on the stage game. Moreover, as we show below, a common case is that unanimity becomes uniquely optimal in the repeated game, thus completely overturning the result of [Proposition 2](#).

Corollary 4. *In the simplified model, when $\alpha \leq 1/2$, there is a range of discount factors for which unanimity ($m = 1$) is uniquely optimal.*

6.3 Consensus

Following [Section 5.4](#), we look for the maximal consensus that can be obtained at the optimal equilibrium of the repeated game. This is achieved by providing to each voter the highest possible incentive (given by the power κ_δ) to vote in favor of decisions that must be accepted at equilibrium. The optimal consensus probability is thus given by

$$P_\delta = \mathbb{P}(u_{[1]} + \kappa_\delta \geq 0) = \mathbb{P}(\theta \geq \alpha - \kappa_\delta) = \frac{1 + \kappa_\delta - \alpha}{2}$$

whenever $\alpha - \kappa_\delta \in [-1, 1]$. This probability can be computed explicitly in the simplified model by plugging in the formulas for κ_δ .

Proposition 3. *The optimal consensus probability P_δ can be decreasing with diversity α .*

[Proposition 3](#) reflects the complex mapping from parameters to equilibrium outcomes in the repeated game. When preferences in the committee are more diverse (α is higher), voting outcomes may be more often unanimous. Hence, the observation of vote tallies may be a poor indicator of preferences within a committee when a norm of cooperation operates. To get some intuition on [Proposition 3](#), note that increasing diversity has two effects. On the one hand, it is more difficult to convince the worst-off individual to vote against her will when diversity is larger. On the other hand, a larger diversity leads to higher intertemporal incentives (i.e. higher κ_δ). When partial cooperation begins, this latter effect is directly proportional to the measure $|e_p|$ of the majority rule's static inefficiency. The result thus asserts that the second effect dominates when the rule is inefficient enough in static terms.

7 Discussion

Malevolent Behavior - Some of the equilibrium payoffs we characterize in [Theorem 1](#) rely on the possibility of punishing individuals for voting in favor of *efficient*

decisions (either rejecting good proposals or accepting bad proposals). One may consider such behavior to be unrealistic, or at the very least undesirable. How much cooperation could be achieved at equilibrium without having to rely on these kind of malevolent incentives? We say that a strategy profile is *non-malevolent* if after any history, implementing the inefficient decision does not lead to a greater continuation utility than implementing the efficient decision. Under the assumption of non-malevolent strategies, the worst equilibrium payoff necessarily coincides with the payoff of the stage-game equilibrium U^0 . The set of equilibrium payoffs is then equal to $[U^0, \bar{V}(\kappa'_\delta)]$, where cutoff $\kappa'_\delta \leq \kappa_\delta$ is given by,

$$\kappa'_\delta = \max \left\{ \kappa \geq 0 \mid (1 - \delta)\kappa = \delta [\bar{V}(\kappa) - U^0] \right\}.$$

The restriction to non-malevolent strategies weakens the optimal punishment (i.e. the lowest equilibrium payoff), thus reducing the largest equilibrium payoff. The set of equilibrium payoffs under non-malevolence is thus always included in the set of equilibrium payoffs characterized in [Theorem 1](#). The inclusion is strict unless there is no cooperation. Yet, we can show that the result on the optimal majority rule in the simplified model ([Theorem 2](#)) remains valid under non-malevolence.

Dropping IIV - If strategies can be conditioned on the detailed history of past votes (as opposed to only past collective decisions), then deviations can be punished individually. In some cases, this may lead to worst equilibrium outcomes than the ones we characterize under IIV. In that sense, our analysis can be understood as providing a lower bound on the maximal equilibrium payoff.

Dropping Symmetry - Our analysis describes the emergence of a norm of cooperation in a stylized setting where committee members are ex-ante symmetric. This is a substantive assumption that allows the precise characterization of equilibrium outcomes ([Theorem 1](#)) and optimal voting rules ([Theorem 2](#)). In many committees, voters are often organized, formally or not, in groups (e.g., political parties in assemblies) which share similar preferences. While the precise analysis seems more difficult in this context, we may speculate that similar cooperation norms

can emerge in two ways. First, a cooperation norm can emerge within groups, voters sometimes voting against their will to better satisfy the overall preference of the group, if a form of symmetry in preferences exists within the group. Second, a cooperation norm can also be sustained across groups, provided that “groups’ preferences” are sufficiently symmetric.

Incomplete information - Our model assumes complete information in the stage game. This assumption is essential for individuals to identify good proposals and punish deviations in order to sustain better outcomes at the equilibrium of the repeated game. Yet, we can show a property of continuity : if the incompleteness of information is low enough, the norm of cooperation that we showcase in our analysis can still emerge in the repeated game.

8 Conclusion

In this paper, we propose a general model of repeated voting in committees and study equilibrium outcomes under alternative majority rules. In contrast with most of the existing literature, which usually focuses on repeated decisions over a single issue under Markovian strategies, we consider a committee making decisions over issues of varying nature, while allowing for history-dependent strategies. Our characterization reveals a complex relationship between the majority rule, the preference distribution, and the set of equilibrium payoffs. We thus turn to a simplified preference domain to perform various comparative statics. Our findings provide theoretical support for two important stylized facts about collective-decision-making in committees: the ubiquity of unanimity as a formal voting rule, and the prevalence of consensus decision-making.

References

Ali, S. N., Bernheim, B. D., Bloedel, A. W., and Battilana, S. C. (2022). Who controls the agenda controls the polity. *arXiv preprint arXiv:2212.01263*.

- Anesi, V. and Seidmann, D. J. (2015). Bargaining in standing committees with an endogenous default. *Review of Economic Studies*, 82:825–867.
- Baron, D. (1996). A dynamic theory of collective goods programs. *The American Political Science Review*, 90(2):316–330.
- Baron, D. and Kalai, E. (1993). The simplest equilibrium of a majority-rule division game. *Journal of Economic Theory*, 61(2):290–301.
- Bernheim, B. D. and Slavov, S. N. (2009). A solution concept for majority rule in dynamic settings. *Review of Economic Studies*, 76:33–62.
- Buchanan, J. and Tullock, G. (1962). The calculus of consent.
- Carrubba, C. J. and Volden, C. (2000). Coalitional politics and logrolling in legislative institutions. *American Journal of Political Science*, 44:261–277.
- Casella, A. and Macé, A. (2021). Does vote trading improve welfare? *Annual Review of Economics*, 13:57–86.
- Casella, A. and Palfrey, T. (2019). Trading votes for votes. a dynamic theory. *Econometrica*, 87:631–652.
- Charroin, L. and Vanberg, C. (2021). Logrolling affects the relative performance of alternative q-majority rules. *Working paper*.
- Duggan, J. and Kalandrakis, A. (2012). Dynamic legislative policy making. *Journal of Economic Theory*, 147(5):1653–1688.
- Dzuida, W. and Loeper, A. (2018). Dynamic pivotal politics. *American Political Science Review*, 112(3):580–601.
- Fischbacher, U. and Schudy, S. (2014). Reciprocity and resistance to comprehensive reform. *Public Choice*, 160(3):411–428.
- Fischbacher, U. and Schudy, S. (2020). Agenda control and reciprocity in sequential voting decisions. *Economic Inquiry*, 58(4):1813–1829.

- Kalandrakis, A. (2004). Why not proportional? *Journal of Economic Theory*, 116:294–322.
- Levin, J. (2003). Relational incentive contracts. *American Economic Review*, 93(3):835–857.
- Maggi, G. and Morelli, M. (2006). Self-enforcing voting in international organizations. *The American Economic Review*, 96(4):1137–1158.
- Mailath, G. J. and Samuelson, L. (2006). Repeated games and reputation - long run relationships.
- Park, R. (1967). The possibility of a social welfare function: Comment. *The American Economic Review*, 57:1300–1304.
- Penn, E. M. (2009). A model of farsighted voting. *American journal of Political Science*, 53(1):36–54.
- Urfalino, P. (2014). The rule of non-opposition: Opening up decision-making by consensus. *Journal of Political Philosophy*.

A Proofs

A.1 Proof of **Theorem 1**

Proof. In the sequel, we denote by H the set of histories in the repeated game.

Definition 1. A payoff $w \in \mathbb{R}$ is decomposable on set $W \subset \mathbb{R}$ if there exists a decision rule $d^w : S \mapsto \{0, 1\}$ and continuation payoffs $r^w : S \mapsto W$ (reward) and $q^w : S \mapsto W$ (punishment) such that:

$$w = \mathbb{E}[(1 - \delta)d^w(\mathbf{u})\bar{u} + \delta r^w(\mathbf{u})] \quad (6)$$

and

$$\forall \mathbf{u} \in S, \quad (1 - \delta)d^w(\mathbf{u})u_p + \delta r^w(\mathbf{u}) \geq (1 - \delta)(1 - d^w(\mathbf{u}))u_p + \delta q^w(\mathbf{u}). \quad (7)$$

The set W is self-generating if any $w \in W$ is decomposable on W .

In words, w is decomposable by (d^w, r^w, q^w) if (i) it is the expected average payoff given by the rule d^w and on-path reward r^w and (ii) for any utility realization $\mathbf{u}$, implementing the action $d^w(\mathbf{u})$ is rational for the pivot given a reward $r^w(\mathbf{u})$ and a punishment $q^w(\mathbf{u})$. For the sequel, it is useful to note that the rationality condition (7) is equivalent to:

$$\forall \mathbf{u} \in S, \quad (1 - \delta)(2d^w(\mathbf{u}) - 1)u_p + \delta(r^w(\mathbf{u}) - q^w(\mathbf{u})) \geq 0 \quad (8)$$

We now state and prove three lemmata before concluding the proof.

Lemma 1. *If W is self-generating, then any $w \in W$ is an equilibrium payoff.*

Proof. This proof is adapted from standard techniques exposed in [Mailath and Samuelson \(2006\)](#). Let W be self-generating and let $w^0 \in W$. Consider the automaton defined by:

- the set of states W
- the initial state w^0
- the output function $f : w \mapsto d^w$
- a transition function $\tau : W \times S \times \{0, 1\} \rightarrow W$, such that

$$\tau(w, \mathbf{u}, d) = \begin{cases} r^w(\mathbf{u}) & \text{if } d = d^w(\mathbf{u}) \\ q^w(\mathbf{u}) & \text{otherwise.} \end{cases}$$

Extend the transition function from $W \times S \times \{0, 1\}$ to $W \times H$, where H denotes the set of (ex-ante) reduced histories, by recursively defining $\tau(w, \emptyset) = w$ and

$$\tau(w, h^t) = \tau(\tau(w, h^{t-1}), \mathbf{u}^t, d^t).$$

Then define the (history-dependent) decision rule d by $d(h, \mathbf{u}) = f(\tau(w^0, h))(\mathbf{u})$ and the reward r and punishment q by $r(h, \mathbf{u}) = r^{\tau(w^0, h)}(\mathbf{u})$ and $q(h, \mathbf{u}) = q^{\tau(w^0, h)}(\mathbf{u})$.

Let $\Delta_i(h, \mathbf{u})$ be the utility difference in state $\tau(w^0, h)$ between following the decision $d(h, \mathbf{u})$ and deviating for an individual with utility u_i (assuming that the individual is pivotal). Formally:

$$\begin{aligned}\Delta_i(h, \mathbf{u}) &= (1 - \delta)d(h, \mathbf{u})u_i + \delta r(h, \mathbf{u}) - \left((1 - \delta)(1 - d(h, \mathbf{u}))u_i + \delta q(h, \mathbf{u}) \right) \\ &= (1 - \delta)(2d(h, \mathbf{u}) - 1)u_i + \delta(r(h, \mathbf{u}) - q(h, \mathbf{u})).\end{aligned}$$

By construction, as W is self-generating, by application of (8), we must have $\Delta_p(h, \mathbf{u}) \geq 0$ for the pivot, for any $h \in H$ and $\mathbf{u} \in S$.

Now, define individual strategy σ_i by:

$$\forall h \in H, \forall \mathbf{u} \in S, \quad \sigma_i(h, \mathbf{u}) = \begin{cases} d(h, \mathbf{u}) & \text{if } \Delta_i(h, \mathbf{u}) \geq 0 \\ 1 - d(h, \mathbf{u}) & \text{otherwise.} \end{cases} \quad (9)$$

Observe first that the strategy profile $\sigma = (\sigma_i)_{i \in N}$ indeed implements the decision rule d :

- if $d(h, \mathbf{u}) = 1$, then for any i with $u_i \geq u_p$, we have $\Delta_i(h, \mathbf{u}) \geq \Delta_p(h, \mathbf{u}) \geq 0$ and thus $\sigma_i(h, \mathbf{u}) = d(h, \mathbf{u}) = 1$. Therefore, the strategy profile σ implements the decision $d = 1 = d(h, \mathbf{u})$ after history h for reform $\mathbf{u}$.
- if $d(h, \mathbf{u}) = 0$, then for any i with $u_i \leq u_p$, we have $\Delta_i(h, \mathbf{u}) \geq \Delta_p(h, \mathbf{u}) \geq 0$ and thus $\sigma_i(h, \mathbf{u}) = d(h, \mathbf{u}) = 0$. Therefore, the strategy profile σ implements the decision $d = 0 = d(h, \mathbf{u})$ after history h for reform $\mathbf{u}$.

Then, as W is self-generating, we have by construction that $U(\sigma) = w^0$. To conclude, let us show that σ is an equilibrium. By definition, the profile σ satisfies Symmetry. It also satisfies Independence of Individual Votes since the transition function τ only depend on the reduced reform $\mathbf{u}$ and the decision d , but not on individual votes (hence $\sigma_i(h, \mathbf{u})$ depends on h only through past reforms and decisions). Finally, by definition of σ_i in (9), it is clear that σ is a subgame-perfect equilibrium of the repeated game that satisfies the criterion of As-If-Pivotal Voting. $\square$

Lemma 2. *The set of equilibrium payoffs E is the largest self-generating set.*

Proof. This proof is adapted from standard techniques exposed in [Mailath and Samuelson \(2006\)](#). By the [Lemma 1](#) above, it suffices to show that the set of equilibrium payoffs E is a self-generating set. For any payoff $w \in E$ with associated equilibrium σ , we let:

- the decision rule d^w be such that $d^w(\mathbf{u})$ is the decision implemented at the first stage under the profile σ for reform $\mathbf{u}$.
- the reward r^w be $r^w(\mathbf{u}) = U(\sigma \mid \mathbf{u}, d^w(\mathbf{u}))$
- the punishment q^w be $q^w(\mathbf{u}) = U(\sigma \mid \mathbf{u}, 1 - d^w(\mathbf{u}))$

As σ is a subgame-perfect equilibrium, the promises r^w and q^w take values in E , as desired. Moreover, we have $w = U(\sigma) = \mathbb{E}[(1 - \delta)d^w(\mathbf{u})\bar{u} + \delta r^w(\mathbf{u})]$ by definition of d^w and r^w . Finally, it is easy to see that [\(7\)](#) is satisfied since σ is an equilibrium:

- if $d^w(\mathbf{u}) = 1$, then there are at least m individuals (where m is the majority rule threshold) for which

$$(1 - \delta)u_i + \delta U(\sigma \mid \mathbf{u}, 1) \geq \delta U(\sigma \mid \mathbf{u}, 0)$$

and thus

$$(1 - \delta)u_p + \delta U(\sigma \mid \mathbf{u}, 1) \geq \delta U(\sigma \mid \mathbf{u}, 0).$$

We thus obtain that [\(7\)](#) is satisfied.

- if $d^w(\mathbf{u}) = 0$, the proof is similar and thus omitted.

□

Lemma 3. *The set of equilibrium payoffs is a compact interval.*

Proof. **Claim 1:** If $W \subset \mathbb{R}$ is a self-generating set such that $U^0 \in W$, then any payoff $w \in (\inf W, \sup W)$ can be decomposed on W .

To show that, it suffices to show that for any payoff w decomposable on W and any $\lambda \in (0, 1)$, the payoff $w' = \lambda w + (1 - \lambda)U^0$ is also decomposable on W . Let w be decomposable by a triplet (d, r, q) on W .

We may write:

$$\begin{aligned}
w &= U^0 + (w - U^0) \\
&= U^0 + \mathbb{E} \left[\underbrace{(1 - \delta)d(\mathbf{u})\bar{u} + \delta r(\mathbf{u}) - (1 - \delta)d^0(\mathbf{u})\bar{u} - \delta U^0}_{z(\mathbf{u})} \right] \\
&= U^0 + \int_S z(\mathbf{u})dF(\mathbf{u}).
\end{aligned}$$

As F has no atom by assumption, there must exist a subdomain $S' \subset S$ such that $\int_{S'} z(\mathbf{u})dF(\mathbf{u}) = \lambda \int_S z(\mathbf{u})dF(\mathbf{u})$. Let us define (d', r', q') by:

$$d'(\mathbf{u}) = \begin{cases} d(\mathbf{u}) & \text{if } \mathbf{u} \in S' \\ d^0(\mathbf{u}) & \text{otherwise} \end{cases}, r'(\mathbf{u}) = \begin{cases} r(\mathbf{u}) & \text{if } \mathbf{u} \in S' \\ U^0 & \text{otherwise} \end{cases}, q'(\mathbf{u}) = \begin{cases} q(\mathbf{u}) & \text{if } \mathbf{u} \in S' \\ U^0 & \text{otherwise} \end{cases}.$$

As $U^0 \in W$, it is clear that both r' and q' take values in W . Moreover, since d^0 is the decision rule implemented by the stage-game equilibrium, it is clear by construction that the rationality condition (7) is satisfied by (d', r', q') . Finally, applying (6), the expected payoff is given by:

$$\begin{aligned}
w' &= \int_{S'} ((1 - \delta)d(\mathbf{u})\bar{u} + \delta r(\mathbf{u})) dF(\mathbf{u}) + \int_{S \setminus S'} ((1 - \delta)d^0(\mathbf{u})\bar{u} + \delta U^0) dF(\mathbf{u}) \\
&= \int_{S'} z(\mathbf{u})dF(\mathbf{u}) + U^0 = \lambda(w - U^0) + U^0 = \lambda w + (1 - \lambda)U^0,
\end{aligned}$$

as desired. Hence $w' = \lambda w + (1 - \lambda)U^0$ is decomposable on W , this concludes the proof of Claim 1.

Claim 2: The set of payoffs generated by a compact interval is a compact interval.

Let $W = [\underline{w}, \bar{w}]$ be a compact interval. Let $\kappa = \frac{\delta}{1 - \delta}(\bar{w} - \underline{w})$. We define $(d^{\max}, r^{\max}, q^{\max})$ by:

$$d^{\max}(\mathbf{u}) = d_{\kappa}^e(\mathbf{u}), \quad r^{\max}(\mathbf{u}) = \bar{w}, \quad q^{\max}(\mathbf{u}) = \underline{w}.$$

In words, $d^{\max}$ is the κ -truncated efficient rule, i.e. it is efficient whenever the pivot stake is below κ , and coincides with the sincere rule otherwise. The reward $r^{\max}$ is always the highest possible and the punishment $q^{\max}$ is always the lowest possible. Clearly, by definition of κ , the triplet $(d^{\max}, r^{\max}, q^{\max})$ satisfies the rationality

condition (7). The payoff induced by the triplet $(d^{\max}, r^{\max}, q^{\max})$ is given by:

$$w^{\max} = (1 - \delta)U(d_{\kappa}^e) + \delta\bar{w}.$$

Now, suppose that there is a triplet (d', r', q') which generates a higher payoff $w' > w^{\max}$. For any $\mathbf{u} \in S$, we have $r'(\mathbf{u}) \leq \bar{w} = r^{\max}(\mathbf{u})$, since W is upper bounded by $\bar{w}$. For any $\mathbf{u} \in S$ with $|u_p| \leq \kappa$, we have $d'(\mathbf{u})\bar{u} \leq d^{\max}(\mathbf{u})\bar{u} = d^e(\mathbf{u})\bar{u}$, since $d^{\max}$ is efficient for such reforms. Thus, to have $w' > w^{\max}$, there must exist $\mathbf{u}$ with $|u_p| > \kappa$ and such that $d'(\mathbf{u})\bar{u} > d^{\max}(\mathbf{u})\bar{u} = d^0(\mathbf{u})\bar{u}$. Suppose without loss of generality that $\bar{u} > 0$. Then we have $d^0(\mathbf{u}) = 0$, and thus $u_p < -\kappa$ and $d'(\mathbf{u}) = 1$. Now, we have:

$$\begin{aligned} (1 - \delta)(2d'(\mathbf{u}) - 1)u_p + \delta(r'(\mathbf{u}) - q'(\mathbf{u})) &\leq (1 - \delta)u_p + \delta(\bar{w} - \underline{w}) \\ &< \delta(\bar{w} - \underline{w}) - (1 - \delta)\kappa = 0. \end{aligned}$$

Thus, the triplet (d', r', q') fails the rationality condition (8). We conclude that $w^{\max}$ is the highest payoff that can be decomposed on W .

Consider now the triplet $(d^{\min}, r^{\min}, q^{\min})$ defined by:

$$d^{\min}(\mathbf{u}) = d_{\kappa}^i(\mathbf{u}), \quad r^{\min}(\mathbf{u}) = \underline{w} + \frac{1 - \delta}{\delta}c(\mathbf{u}, d_{\kappa}^i(\mathbf{u})), \quad q^{\min}(\mathbf{u}) = \underline{w}.$$

In words, $d^{\min}$ is the κ -truncated cost-adjusted inefficient rule, i.e. it is inefficient whenever the pivot stake is below κ and the pivot utility is of the same sign as the average utility but with a lower magnitude, and it coincides with the sincere rule otherwise. The reward $r^{\min}$ is the minimum needed to provide incentives to the pivot to vote against her will when the prescribed decision is inefficient and insincere, it is the lowest possible payoff otherwise. The punishment $q^{\min}$ is always the lowest possible.

Clearly, by construction, the triplet $(d^{\min}, r^{\min}, q^{\min})$ satisfies the rationality condition (7). Now, suppose that there is a triplet (d', r', q') which generates a lower payoff $w' < w^{\min}$, where $w^{\min}$ is the expected payoff generated by $(d^{\min}, r^{\min}, q^{\min})$,

that is:

$$w^{\min} = \mathbb{E}[(1 - \delta)d^{\min}(\mathbf{u})\bar{u} + \delta r^{\min}(\mathbf{u})] = (1 - \delta)(U(d_{\kappa}^i) + C(d_{\kappa}^i)) + \delta \underline{w}.$$

Then, there must exist $\mathbf{u} \in S$ such that the following condition holds:

$$(1 - \delta)d'(\mathbf{u})\bar{u} + \delta r'(\mathbf{u}) < (1 - \delta)d^{\min}(\mathbf{u})\bar{u} + \delta r^{\min}(\mathbf{u}). \quad (10)$$

Condition (10) can only be satisfied if at least one of the two following inequalities hold: $d^{\min}(\mathbf{u}) = d^e(\mathbf{u})$ or $r^{\min}(\mathbf{u}) > \underline{w}$. We thus consider two cases:

- if $r^{\min}(\mathbf{u}) > \underline{w}$, then by construction $r^{\min}$, we must have $c(\mathbf{u}, d_{\kappa}^i(\mathbf{u})) > 0$. We thus have $|u_p| < \kappa$ and (abusing notation, as u_p and $\bar{u}$ can be either both positive or both negative) we have $u_p \in (0, \bar{u})$. Under these conditions, we know that $d^{\min}(\mathbf{u}) = 1 - d^e(\mathbf{u}) = 1 - d^0(\mathbf{u})$. It follows that $d'(\mathbf{u})\bar{u} \geq d^{\min}(\mathbf{u})\bar{u}$, so that condition (10) implies that $r'(\mathbf{u}) < r^{\min}(\mathbf{u}) = \underline{w} + \frac{1-\delta}{\delta}|u_p|$. There are two subcases to consider:
 - if $d'(\mathbf{u}) = d^{\min}(\mathbf{u}) = 1 - d^0(\mathbf{u})$, we obtain a contradiction by observing that, since $r'(\mathbf{u}) < r^{\min}(\mathbf{u})$ and $q'(\mathbf{u}) \geq \underline{w} = q^{\min}(\mathbf{u})$, we must have $r'(\mathbf{u}) - q'(\mathbf{u}) < r^{\min}(\mathbf{u}) - q^{\min}(\mathbf{u})$. Therefore, the triplet (d', r', q') cannot satisfy the rationality condition (7) at $\mathbf{u}$ (recall that by construction, $r^{\min}(\mathbf{u}) - q^{\min}(\mathbf{u})$ is the smallest difference between continuations payoffs that can incentivize the pivot individual to deviate from sincere voting).
 - if $d'(\mathbf{u}) = d^0(\mathbf{u}) = d^e(\mathbf{u})$, we obtain a contradiction with (10), as since $u_p \in (0, \bar{u})$, we have:

$$\begin{aligned} (1 - \delta)d'(\mathbf{u})\bar{u} + \delta r'(\mathbf{u}) &\geq (1 - \delta)d^e(\mathbf{u})\bar{u} + \delta \underline{w} \\ &> (1 - \delta)((1 - d^e(\mathbf{u}))\bar{u} + |u_p|) + \delta \underline{w} \\ &> (1 - \delta)d^{\min}(\mathbf{u})\bar{u} + \delta r^{\min}(\mathbf{u}). \end{aligned}$$

- if $d^{\min}(\mathbf{u}) = d^e(\mathbf{u})$, then we have (by construction of the truncated cost-adjusted inefficient rule $d^{\min}$) $c(\mathbf{u}, d^{\min}(\mathbf{u})) = 0$ and thus $r^{\min}(\mathbf{u}) = \underline{w}$ and $d^0(\mathbf{u}) = d^e(\mathbf{u})$. For (10) to hold, as $r'(\mathbf{u}) \geq \underline{w}$, we must have $d'(\mathbf{u}) = 1 -$

$d^e(\mathbf{u}) = 1 - d^0(\mathbf{u})$. We consider two subcases:

- if $|u_p| \leq \kappa$, then we also have (by construction of the truncated cost-adjusted inefficient rule $d^{\min}$) that $|u_p| > |\bar{u}|$. For the triplet (d', r', q') to satisfy the rationality condition (7), we must have $r'(\mathbf{u}) - q'(\mathbf{u}) \geq \frac{1-\delta}{\delta}|u_p|$, and thus $r'(\mathbf{u}) \geq \underline{w} + \frac{1-\delta}{\delta}|u_p|$. We obtain a contradiction with (10), as since $|u_p| > |\bar{u}|$, we have:

$$\begin{aligned} (1 - \delta)d'(\mathbf{u})\bar{u} + \delta r'(\mathbf{u}) &\geq (1 - \delta)((1 - d^e(\mathbf{u}))\bar{u} + |u_p|) + \delta \underline{w} \\ &> (1 - \delta)d^e(\mathbf{u})\bar{u} + \delta \underline{w} \\ &> (1 - \delta)d^{\min}(\mathbf{u})\bar{u} + \delta r^{\min}(\mathbf{u}). \end{aligned}$$

- if $|u_p| > \kappa$, then we obtain a contradiction by observing that, since $r'(\mathbf{u}) - q'(\mathbf{u}) \leq \bar{w} - \underline{w} = \frac{1-\delta}{\delta}\kappa$, the triplet (d', r', q') cannot satisfy the rationality condition (7) at $\mathbf{u}$.

We conclude that $w^{\min}$ is the lowest payoff that can be decomposed on W .

Claim 3: If a bounded interval I is self-generating, then its closure $\bar{I}$ is also self-generating.

If I is self-generating, then any payoff in I is decomposable on I , and thus also on $\bar{I}$. The set of payoffs that can be decomposed on $\bar{I}$ is thus a compact interval (by Claim 2) which contains I . Thus any payoff in $\bar{I}$ can be decomposed on $\bar{I}$, i.e. $\bar{I}$ is self-generating.

To conclude, we know that the set of equilibrium payoffs E contains U^0 (repeating the stage-game equilibrium is an equilibrium of the repeated game). As E is self-generating, it must be a (bounded) interval by Claim 1. As E is the largest self-generating set by the previous lemma, it must be closed by Claim 3, and thus compact. This concludes the proof of Lemma 3. $\square$

We are now equipped to finish the proof of Theorem 1. We know from the previous lemma that $E = [\underline{w}, \bar{w}]$. Let $w^{\max}$ the maximal payoff that can be decomposed on E . As E is self-generating, we have that $w^{\max} \geq \bar{w}$. Now, if $w^{\max} > \bar{w}$, then this would contradict the fact that E must be the largest self-generating set.

Hence, we must have $w^{\max} = \bar{w}$. Similarly, if we note $w^{\min}$ the minimal payoff that can be decomposed on E , we have $w^{\min} = \underline{w}$.

The formula for $w^{\max}$ in the proof is $w^{\max} = (1 - \delta)U(d_\kappa^e) + \delta\bar{w}$. We thus obtain $w^{\max} = U(d_\kappa^e) = \bar{V}(\kappa)$. The formula for $w^{\min}$ in the proof is $w^{\min} = (1 - \delta)(U(d_\kappa^i) + C(d_\kappa^i)) + \delta\underline{w}$. We thus obtain $w^{\min} = U(d_\kappa^i) + C(d_\kappa^i) = \underline{V}(\kappa)$.

We thus obtained that $E = [\underline{V}(\kappa), \bar{V}(\kappa)]$ for $\kappa \geq 0$ such that $(1 - \delta)\kappa = \delta(\bar{V}(\kappa) - \underline{V}(\kappa))$. To finish, suppose that there is $\kappa' > \kappa$ such that $(1 - \delta)\kappa' = \delta(\bar{V}(\kappa') - \underline{V}(\kappa'))$. Following the arguments in the proof of the previous lemma, we obtain that $[\underline{V}(\kappa'), \bar{V}(\kappa')]$ is self-generating, but this set strictly contains E , this provides a contradiction with the fact that E is the largest self-generating set. This concludes the proof. $\square$

A.2 Proof of Proposition 1

Proof. First, it is straightforward that the highest consensus will be achieved at an optimal equilibrium if, for any reform $\mathbf{u}$ with prescribed decision $d_{\kappa_\delta}^e(\mathbf{u}) = 1$, the reward will be the highest possible, equal to $\bar{w}_\delta$, and the punishment will be the lowest possible, equal to $\underline{w}_\delta$. Thus, such an optimal equilibrium will be stationary, implementing a stage-game profile $v = (v_i)_{i \in N}$ with $v_i(\mathbf{u}) = \mathbb{1}(\{u_i \geq -\kappa_\delta\})$ whenever $d_{\kappa_\delta}^e(\mathbf{u}) = 1$.

As mentioned in the text below Proposition 1, it is straightforward that more good reforms will be adopted and that each such reform is more likely to receive consensus when δ increases. Therefore, to show that P_δ is increasing, it suffices to show that a bad reform never receives consensus (even if approved).

Let $\mathbf{u}$ be a bad reform, i.e. $\bar{u} \leq 0$ and assume that it is approved at an optimal equilibrium, i.e. $d_{\kappa_\delta}^e(\mathbf{u}) = 1$. Then, by definition of the truncated efficient rule, it must be that $u_p > \kappa_\delta$. Now observe that we can upper bound the lowest utility $u_{[1]}$ by lower bounding the average utility

$$0 \geq \bar{u} > mn\kappa_\delta + (1 - m)nu_{[1]}.$$

We thus have $u_{[1]} < -\frac{m}{1-m}\kappa_\delta \leq \kappa_\delta$, since $m \geq 1/2$. It follows that reform $\mathbf{u}$ cannot be unanimously approved at an optimal equilibrium.

We have thus shown that P_δ is increasing. As the support S of reform distribution F is bounded, while the function κ_δ is increasing and unbounded (since $\bar{w}_\delta > U^0$ for $\delta > \delta^F$), there must exist $\delta^C \in (0, 1)$ above which we have $\kappa_\delta \geq -\min_{\mathbf{u} \in S} u_{[1]}$. Then, consensus is maximal and any good reform is adopted with consensus.

□

A.3 Proof of Theorem 2

Proof. The first step of the proof consists in providing closed-form formulas and comparative results for all variables of interest (b_δ , $\bar{w}_\delta$, discount factor thresholds) for any possible value of discount factor δ and inefficiency gap Δ (throughout the proof we abuse notation and write Δ for Δ^e). It will be useful to divide our inquiry in two cases, corresponding to two regions of values for Δ , that is $\Delta \in (0, 1/2]$, and then $\Delta \in (1/2, 1]$.

A.3.1 Preliminaries

Note that for any $\kappa \in [0, \Delta]$, we may write

$$\bar{V}(\kappa) = \int_{u_c \geq \kappa} \bar{u} dF(\mathbf{u}) = \int_{\theta \geq \Delta - \kappa} \theta \frac{d\theta}{2} = \frac{1 - (\kappa - \Delta)^2}{4}.$$

Hence, we obtain the general formula

$$\bar{V}(\kappa) = \begin{cases} \frac{1 - (\kappa - \Delta)^2}{4} & \text{if } 0 \leq \kappa \leq \Delta \\ \frac{1}{4} & \text{if } \kappa \geq \Delta. \end{cases} \quad (11)$$

Note that when $\Delta \leq 1 - \kappa$, we have

$$\underline{V}(\kappa) = \int_{\Delta}^{\Delta + \kappa} (\theta - \Delta) \frac{d\theta}{2} + \int_{\Delta + \kappa}^1 \theta \frac{d\theta}{2} = \frac{1 - \Delta(\Delta + 2\kappa)}{4},$$

while when $1 - \kappa \leq \Delta \leq 1$, we have

$$\underline{V}(\kappa) = \int_{\Delta}^1 (\theta - \Delta) \frac{d\theta}{2} = \frac{(1 - \Delta)^2}{4}.$$

Hence, we obtain the general formula

$$\underline{V}(\kappa) = \begin{cases} \frac{1 - \Delta(\Delta + 2\kappa)}{4} & \text{if } \kappa \leq 1 - \Delta \\ \frac{(1 - \Delta)^2}{4} & \text{if } \kappa \geq 1 - \Delta. \end{cases} \quad (12)$$

A.3.2 Equilibrium Payoffs

We separate our analysis between Case 1, $\Delta \in (0, 1/2]$, and Case 2, $\Delta \in (1/2, 1]$.

Case 1: $\Delta \in (0, 1/2]$

In that case, we have $\Delta \leq 1 - \Delta$. By application of (11) and (17), we obtain

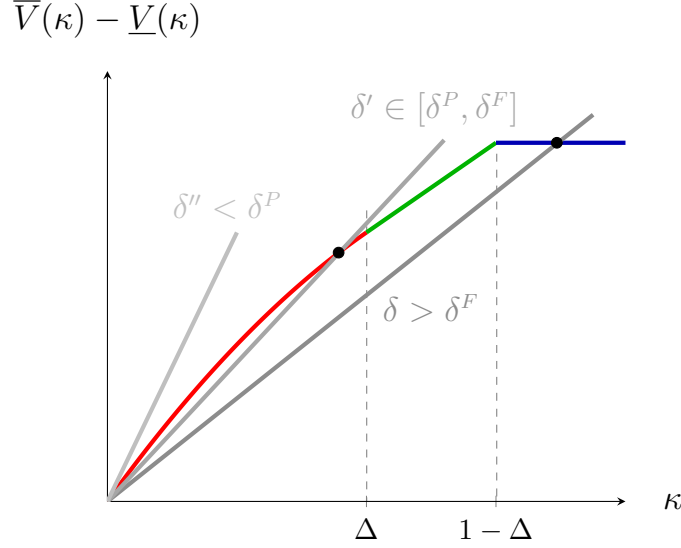
$$\bar{V}(\kappa) - \underline{V}(\kappa) = \begin{cases} \kappa \left(\Delta - \frac{\kappa}{4} \right) & \text{if } \kappa \leq \Delta \\ \frac{\Delta(\Delta + 2\kappa)}{4} & \text{if } \Delta \leq \kappa \leq 1 - \Delta \\ \frac{\Delta(2 - \Delta)}{4} & \text{if } \kappa \geq 1 - \Delta. \end{cases} \quad (13)$$

Observe that the function $\bar{V}(\kappa) - \underline{V}(\kappa)$ is strictly concave for $\kappa \in [0, \Delta]$, then affine for $\kappa \in [\Delta, 1 - \Delta]$, and finally constant for $\kappa \in [1 - \Delta, +\infty)$. Moreover, its derivative is continuous at $k = \Delta$. Equation (4) thus admits no solution (if δ is too low) or a unique solution (if δ is high enough).

The determination of κ_{δ} in Case 1 is illustrated in Figure 3 below. The colored curve corresponds to the difference $\bar{V}(\kappa) - \underline{V}(\kappa)$, while the grey lines represent functions $\kappa \mapsto (1 - \delta)\kappa/\delta$, for different values of the discount factor δ . As κ increases, both the highest and the lowest sustainable payoff increase. Once κ reaches Δ , the highest sustainable payoff reaches its maximum $U^e = 1/4$, while the lowest sustainable payoff keeps decreasing (green part). For κ larger than $1 - \Delta$,

the lowest sustainable payoff reaches its minimum $(1 - \Delta)^2/4$ (blue part).

Figure 3: Optimal Equilibrium in Case 1.



Let thresholds δ^P and δ^F such that,

$$\frac{1 - \delta^P}{\delta^P} = \frac{\partial \left(\bar{V}(\kappa) - \underline{V}(\kappa) \right)}{\partial \kappa} \Big|_{\kappa=0} = \Delta, \quad \frac{1 - \delta^F}{\delta^F} = \frac{\left(\bar{V}(\Delta) - \underline{V}(\Delta) \right)}{\Delta} = \frac{3\Delta}{4}. \quad (14)$$

We identify three regimes of cooperation as a function of δ .

Regime 1: $\delta \leq \delta^P$. The grey line lies above the colored curve for any positive value of κ , and we get that $\kappa_\delta = 0$; there is no cooperation and the unique equilibrium payoff is U^0 .

Regime 2: $\delta \in (\delta^P, \delta^F)$. The grey line intersects the colored curve for a value of $\kappa_\delta \in (0, \Delta)$ (red part) defined by,

$$\kappa_\delta = 4 \left(\Delta - \frac{1 - \delta}{\delta} \right). \quad (15)$$

and the associated highest equilibrium payoff is given by,

$$\bar{w}_\delta = U^e - 4 \left(\frac{1}{\delta} - \frac{1}{\delta^F} \right)^2. \quad (16)$$

We then get a regime of partial cooperation as $\bar{w}_\delta \in (U^0, U^e)$.

Regime 3: $\delta \geq \delta^F$. The grey line intersects either the green or the blue line for a value of κ greater than Δ . In that case, we get full cooperation as $\bar{w}_\delta = U^e$.

Comparative Statics: κ_δ is increasing in Δ in regime 2 and thus $\bar{w}_\delta$ is also increasing in Δ in regime 2. Both thresholds δ^P and δ^F are decreasing in Δ .

Case 2: $\Delta \in (1/2, 1]$

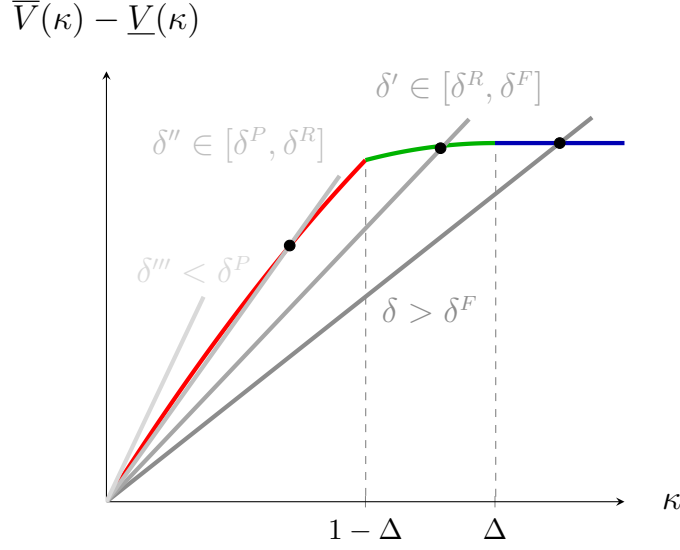
In that case, we have $\Delta \geq 1 - \Delta$. By application of (11) and (17), we obtain

$$\bar{V}(\kappa) - \underline{V}(\kappa) = \begin{cases} \kappa \left(\Delta - \frac{\kappa}{4} \right) & \text{if } \kappa \leq 1 - \Delta \\ \frac{-\kappa^2 + 2\Delta(1 + \kappa) - 2\Delta^2}{4} & \text{if } 1 - \Delta \leq \kappa \leq \Delta \\ \frac{\Delta(2 - \Delta)}{4} & \text{if } \kappa \geq \Delta. \end{cases} \quad (17)$$

Observe that the function $\bar{V}(\kappa) - \underline{V}(\kappa)$ is strictly concave for $\kappa \in [0, 1 - \Delta]$, then strictly concave for $\kappa \in [1 - \Delta, \Delta]$, and finally constant for $\kappa \in [\Delta, +\infty)$. Moreover, its derivative is (discontinuously) decreasing at $k = 1 - \Delta$. Equation (4) thus admits no solution (if δ is too low) or a unique solution (if δ is high enough).

The determination of κ_δ in Case 2 is illustrated in Figure 4 below. As κ increases, both the highest and the lowest sustainable payoff increase. Once κ reaches $1 - \Delta$, the lowest sustainable payoff reaches its minimum $(1 - \Delta)^2/4$, while the highest sustainable payoff keeps increasing (green part). For κ larger than Δ , the highest sustainable payoff reaches its maximum $1/4$ (blue part).

Figure 4: Optimal Equilibrium in Case 2.



Let threshold δ^P as in (14), and define new thresholds δ^R and δ^F such that,

$$\frac{1 - \delta^R}{\delta^R} = \frac{(\bar{V}(1 - \Delta) - \underline{V}(1 - \Delta))}{\Delta} = \frac{5\Delta - 1}{4}, \quad \frac{1 - \delta^F}{\delta^F} = \frac{(\bar{V}(\Delta) - \underline{V}(\Delta))}{\Delta} = \frac{2 - \Delta}{4}$$

We identify four regimes of cooperation as a function of δ .

Regime 1: $\delta \leq \delta^P$. The grey line lies above the colored curve for any positive value of κ , and we get that $\kappa_\delta = 0$.

Regime 2: $\delta \in (\delta^P, \delta^R)$. The grey line intersects the colored curve for a value $\kappa_\delta < 1 - \Delta$ (red part) defined by equation (15), as in Case 1. The highest equilibrium payoff $\bar{w}_\delta$ is given by formula (16).

Regime 3: $\delta \in (\delta^R, \delta^F)$. The grey line intersects the colored curve for a value $\kappa_\delta \in (1 - \Delta, \Delta]$ (green part) defined by,

$$(1 - \delta)\kappa_\delta = \delta \left(\frac{-\kappa^2 + 2\Delta(1 + \kappa) - 2\Delta^2}{4} \right).$$

Simplifying, we obtain the quadratic equation,

$$\kappa^2 + 4 \left(\frac{1-\delta}{\delta} - \frac{\Delta}{2} \right) \kappa + 2\Delta(\Delta - 1) = 0. \quad (18)$$

Since $\Delta - 1 < 0$, equation (18) admits a unique positive solution, given by,

$$\kappa_\delta = \frac{1}{2} \left(\frac{1-\delta}{\delta} - \frac{\Delta}{2} \right) \left(-1 + \sqrt{1 + \frac{\Delta(1-\Delta)}{2 \left(\frac{1-\delta}{\delta} - \frac{\Delta}{2} \right)^2}} \right).$$

The highest equilibrium payoff is now given by,

$$\bar{w}_\delta = \frac{1 - (\Delta - \kappa_\delta)^2}{4}.$$

Regime 4: $\delta \geq \delta^F$. The grey line intersects the colored curve for a value of κ greater than Δ (blue part). In that case, we get full cooperation as $\bar{w}_\delta = U^e$.

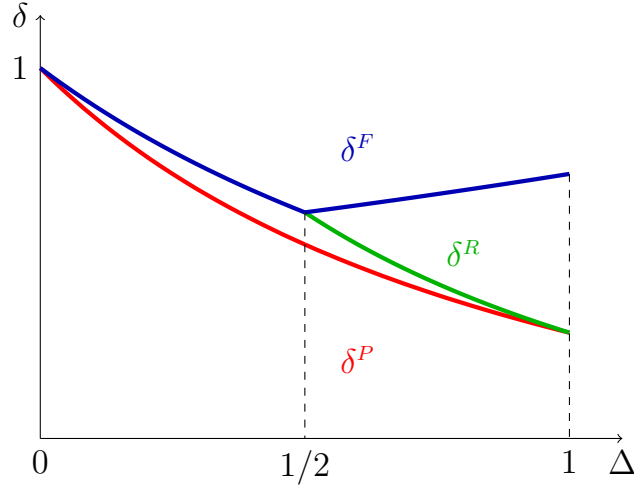
Comparative Statics: thresholds δ^P and δ^R are decreasing with respect to Δ , while threshold δ^F is increasing with respect to Δ . As long as $\delta < \delta^R$ (regimes 1 and 2), κ_δ , and $\bar{w}_\delta$ are increasing in Δ , as in Case 1. By contrast, when $\delta \in (\delta^R, \delta^F)$ (regime 3) $\bar{w}_\delta$ is decreasing in Δ . Indeed, by differentiating (18), we obtain that,

$$\frac{\partial(\Delta - \kappa)}{\partial\Delta} = \frac{\Delta + 2 \left(\frac{1-\delta}{\delta} \right) - 1}{\kappa - \Delta + 2 \left(\frac{1-\delta}{\delta} \right)}.$$

The numerator can be shown to be positive as $\delta \leq \delta^F$ implies that $\frac{1-\delta}{\delta} \geq 1 - \frac{\Delta}{2}$. The denominator is also positive by definition of κ_δ in (18).

Figure 5 below represents the various discount thresholds as a function of the inefficiency gap Δ . We note that full cooperation is most easily achieved when $\Delta = 1/2$.

Figure 5: Cooperation Thresholds.



A.3.3 Optimal voting rule(s).

We now characterize the set of optimal voting rules, i.e. the voting rules which maximize the highest equilibrium payoff $\bar{w}_\delta$. First, observe that $\bar{w}_\delta$ is increasing in δ , and,

- constant in Δ for $\delta < \delta^P$,
- increasing in Δ for $\delta^P < \delta < \delta^F$ when $\Delta \leq 1/2$, and for $\delta^P < \delta < \delta^R$ when $\Delta \geq 1/2$, and
- decreasing in Δ for $\delta^R < \delta < \delta^F$ when $\Delta \geq 1/2$.

Hence, as long as full cooperation does not obtain for any Δ , i.e. $\delta < \delta^F(\Delta = 1/2)$, we get that:

- the best rule among those that do not allow for partial cooperation is simple majority (by application of [Proposition 2](#)).
- the best rule among those that allow for partial cooperation is the highest Δ that remains below $(\delta^R)^{-1}$.

Moreover, the payoff of the first rule is constant and strictly smaller than U^e , while the payoff of the second one increases with δ , converging to U^e as δ goes to $\min_\Delta \delta^F(\Delta)$. We conclude that there is a threshold $\underline{\delta}$ above which the second rule

is optimal. This rule can be written $m^*(\delta)$ and it must be (weakly) decreasing in δ since $(\delta^R)^{-1}$ is itself decreasing. This concludes the proof.

□

A.4 Proof of **Proposition 3**

Proof. Consider the optimal degree of cooperation κ_δ in the first regime of cases 1 and 2 in the previous proof of **Theorem 2**. In this regime, $\kappa_\delta = 4(\Delta - \frac{1-\delta}{\delta})$. We thus have $\frac{\partial \kappa_\delta}{\partial \alpha} = 4|e_p|$. Hence, we obtain that $\frac{\partial P_\delta}{\partial \alpha} > 0$ whenever $4|e_p| > 1$, i.e. $|e_p| > 1/4$. □