

Anesi, Vincent; Bowen, T. Renee

Working Paper

Policy experimentation, redistribution and voting rules

CeDEx Discussion Paper Series, No. 2018-09

Provided in Cooperation with:

The University of Nottingham, Centre for Decision Research and Experimental Economics (CeDEx)

Suggested Citation: Anesi, Vincent; Bowen, T. Renee (2018) : Policy experimentation, redistribution and voting rules, CeDEx Discussion Paper Series, No. 2018-09, The University of Nottingham, Centre for Decision Research and Experimental Economics (CeDEx), Nottingham

This Version is available at:

<https://hdl.handle.net/10419/200428>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.



CENTRE FOR DECISION RESEARCH & EXPERIMENTAL ECONOMICS



The University of
Nottingham

UNITED KINGDOM • CHINA • MALAYSIA

Discussion Paper No. 2018-09

Vincent Anesi and
T. Renee Bowen

August 2018

**Policy Experimentation,
Redistribution and Voting
Rules**

CeDEx Discussion Paper Series
ISSN 1749 - 3293



CENTRE FOR DECISION RESEARCH & EXPERIMENTAL ECONOMICS

The Centre for Decision Research and Experimental Economics was founded in 2000, and is based in the School of Economics at the University of Nottingham.

The focus for the Centre is research into individual and strategic decision-making using a combination of theoretical and experimental methods. On the theory side, members of the Centre investigate individual choice under uncertainty, cooperative and non-cooperative game theory, as well as theories of psychology, bounded rationality and evolutionary game theory. Members of the Centre have applied experimental methods in the fields of public economics, individual choice under risk and uncertainty, strategic interaction, and the performance of auctions, markets and other economic institutions. Much of the Centre's research involves collaborative projects with researchers from other departments in the UK and overseas.

Please visit <http://www.nottingham.ac.uk/cedex> for more information about the Centre or contact

Suzanne Robey
Centre for Decision Research and Experimental Economics
School of Economics
University of Nottingham
University Park
Nottingham
NG7 2RD
Tel: +44 (0)115 95 14763
suzanne.robey@nottingham.ac.uk

The full list of CeDEX Discussion Papers is available at

<http://www.nottingham.ac.uk/cedex/publications/discussion-papers/index.aspx>

Policy Experimentation, Redistribution and Voting Rules*

Vincent Anesi[†]
T. Renee Bowen[‡]

August 8, 2018

Abstract

We study conditions under which optimal policy experimentation can be implemented by a committee. We consider a dynamic bargaining game in which committee members choose to implement either a risky reform or a safe alternative with known returns each period. We find that when no redistribution is allowed the unique equilibrium outcome is generically inefficient. When committee members are allowed to redistribute resources (even arbitrarily small amounts), there always exists an equilibrium that supports optimal experimentation for any voting rule with no veto players. With veto players, however, optimal policy experimentation is possible only with a sufficient amount of redistribution. We conclude that veto rights are more of an obstacle to optimal policy experimentation than constraints on redistribution.

*We thank Nageeb Ali, David Baron, Laura Doval, Wiola Dziuda, John Ferejohn, Navin Kartik, Lewis Kornhauser, Facundo Piguillem, Daniel Seidmann, Ken Shotts, Joel Watson, and participants in the Stanford Political Economy Theory seminar, the Summer Workshop in Political Economy at Stony Brook University, the UCSD theory mini conference, the NYU Law and Economics Colloquium, the UEA Economic Theory Workshop, the UCSD theory seminar, the Africa meetings of the Econometric Society, for helpful comments and suggestions. We also thank Wenlong Zhao for excellent research assistance.

[†]University of Nottingham

[‡]University of California, San Diego and NBER

1 Introduction

Committees often must choose to embark on risky policy experiments. Examples are governments implementing reforms, and international organizations agreeing to new policies. One objective of such a committee is to embark on a policy experiment if the expected aggregate benefits exceed the aggregate costs, and, conversely, quit an existing policy experiment when the likelihood that the benefits will materialize becomes sufficiently small. However, the decision to begin or end a policy experiment is made by individual committee members acting in their own interests. In spite of the promise of aggregate benefits, if a sufficient number of members do not directly gain, then policy experiments may fail to be implemented. Similarly, if a sufficiently large coalition of vested interests supports a failed experiment, it may persist at the expense of other members. The political economy literature on policy experimentation — which we elaborate on below — has so far developed independently of that on redistribution. However, redistribution might be a critical factor to circumvent the vested interests opposing a socially beneficial reform or the cessation of a failed reform: the gains to winners should be redistributed in such a way that losers are fully compensated and the right level of experimentation occurs. Such compensatory transfers are difficult to implement in practice because, as Acemoglu et al. (2015) points out, there can be a number of (exogenous) direct constraints on redistribution. Even if some redistribution is allowed, committees may face another critical constraint: potential winners from a policy experiment may not be able to *commit* to compensating the losers.

The observations above motivate the three main questions addressed in this paper. First, under what circumstances, if any, can committees achieve socially efficient policy experimentation without redistribution? Second, if efficient experimentation cannot be achieved without redistribution, can it be achieved if benefits can be redistributed among members? Is there some minimal level of redistribution necessary for this to happen? And third, how do the answers depend on voting rules?

To answer these questions we present a dynamic committee bargaining model that combines policy experimentation and redistribution. A committee (such as a legislature, the members of an international body, or a central bank governing board) meets each period

to decide policy. Policy has three components: the choice to implement a risky reform or revert to a safe (known) alternative, a tax rate, and the choice of how to distribute tax revenues. We model the constraint on redistribution as an exogenous maximum tax rate. We assume that this is chosen separately from reforms and redistribution. For example, U.S. tax legislation is primarily driven by the U.S. Treasury and the office of the President, whereas other policies related to reforms and redistribution may come from the House of Representatives or the Senate. Other possible interpretations are constitutional constraints and the threat of capital flight (see Acemoglu et al., 2015), or that only part of individual benefits are observable or transferable.

Policy is assumed to continue unless changed by the committee, and, in this sense, exhibits an *endogenous status quo* feature.¹ More precisely, in each period, policy making is governed by the following protocol adapted from Diermeier et al. (2017). Committee members have the opportunity to propose amendments to the ongoing status quo in random order, determined at the start of the period. If a proposal is voted up, then it is implemented in that period and becomes the next period's status quo. If all proposers fail to amend the status quo, then it is implemented and remains in place until the next period.

The choice between a risky reform and a safe alternative is modeled as a bandit problem in the spirit of Keller et al. (2005), but with heterogeneous payoffs across committee members. The safe alternative generates a certain benefit if selected, but these benefits differ across members. If the reform is good, then it generates benefits stochastically. These benefits also differ across committee members. If the reform is not good, then it never generates benefits. There is a prior belief that the reform is good and this is common to all committee members. With each failed attempt at reform, all committee members update their beliefs about whether the reform is good or not according to Bayes' rule. Thus, with each failure the belief that the reform is good decreases, and the expected payoff from the

¹This endogenous status quo feature of dynamic policymaking has been used in a number of recent papers and in a variety of policy settings. These include Kalandrakis (2004) (pure redistribution), Bowen et al. (2014) (entitlement programs), Piguillem and Riboni (2015) (public spending), Dziuda and Loeper (2016) (binary policy).

reform also decreases. If the reform produces a success, then all committee members know the reform is good with probability one. In the absence of bargaining, a utilitarian social planner follows a rule according to which the reform is attempted until the belief that it is good is sufficiently small. We call this the *optimal stopping rule*.

We begin the analysis of the dynamic bargaining game with the case in which no redistribution is allowed and find that the committee, in general, does not implement the optimal stopping rule. More precisely, we show that there is a unique equilibrium outcome and it is inefficient. Either experimentation never occurs, or the committee implements a stopping rule whose cutoff belief for ending experimentation is the lowest median among the committee members' ideal cutoffs. We conclude that whenever this cutoff differs from the social planner's ideal cutoff, then the committee fails to implement the optimal stopping rule.

We then turn to the other extreme case, in which the committee can freely redistribute revenues among its members. Perhaps unsurprisingly, we find that all equilibria sustain the optimal stopping rule. The heterogeneity of the relative benefits from experimentation among committee members is immaterial when those benefits can be fully redistributed. Coupled with the no-redistribution case, this result seems to suggest that constraints on redistribution are the main impediments to socially efficient experimentation. But our analysis of the more realistic case in which redistribution is limited reveals that this is not the complete picture. Indeed, with a non-collegial voting rule (i.e., without veto players), socially efficient experimentation is attainable as long as the exogenous constraint on redistribution is relaxed (even marginally). In contrast, with collegial rules, a minimum level of redistribution must be permitted to sustain the optimal stopping rule.

We conclude from our analysis that in situations where redistribution is limited, veto rights, not the constraints on redistribution, are the main obstacle to efficient policy experimentation. Committees with veto players, who can block proposals to move away from current policies, constitute an empirically important class of institutions (e.g., the United Nations Security Council and presidential veto power in the US Congress). Our focus on experimentation and redistribution thus offers a different perspective to exist-

ing analyses of the normative implications of voting rules on committee decision making, which have mainly concentrated on the information-aggregation channel. In our collective-experimentation framework, non-collegial voting rules permit socially efficient experimentation for a larger set of parameters than do collegial rules, so that the former rules “dominate” the latter in the sense of Bouton et al. (2018).

Related literature This project is most closely related to the literature on collective experimentation and voting rules, including Strulovici (2010), and Messner and Polborn (2012). Like our paper, these papers study how various voting rules affect incentives to experiment in committees. Messner and Polborn (2012) consider a two-period model and compare the optimality of different majority rules. We consider an infinite horizon model with general decision rules (collegial and non-collegial), which include all non-unanimity majority rules as examples of non-collegial, and unanimity as an example of collegial. In an infinite horizon setting, Strulovici (2010) shows that efficient policy experimentation cannot be sustained with voting. This occurs as voters learn whether or not the policy will be beneficial to them. In contrast to Strulovici (2010), we assume that agents know their potential future benefit from experimentation, but there is a common uncertainty about whether the reform is good. In this sense, a conflict exists between voters prior to beginning experimentation and continues throughout experimentation. We show that this conflict can be mitigated with a sufficient level of redistribution. Acemoglu et al. (2015) consider a model of collective experimentation over political institutions, in which trials of the risky alternative involve uncertainty not only about its payoff implications but also about who controls political power. Strulovici (2010), Messner and Polborn (2012) and Acemoglu et al. (2015) do not consider redistribution. Other papers considering policy experimentation and politics include Majumdar and Mukand (2004), Volden et al. (2008), Cai and Treisman (2009), Callander (2011), Callander and Hummel (2014), Hirsch (2016) and Freer et al. (2018). These papers do not consider dynamic committee bargaining and policy experimentation.

Our paper compares committees with non-collegial voting rules (without veto players),

to committees with collegial rules (with veto players). This focus connects the current paper to a rich literature on veto players. This literature includes the seminal works by Tsebelis (2002) and Krehbiel (1998). As with a number of papers in this literature, we conclude that the presence of veto players can generate inefficiencies.² The inefficiency we identify is the inability of a committee to optimally experiment with reform when this requires the decision to both optimally begin the experiment and end the experiment when it proves to be unsuccessful.³

We consider that policies, once implemented, can only be changed with a new round of voting, and hence what we do relates to the literature on bargaining with an endogenous status quo, pioneered by Baron (1996). Our framework extends the stationary divide-the-dollar framework studied in Kalandrakis (2004, 2010), Battaglini and Palfrey (2012), Bowen and Zahran (2012), Baron and Bowen (2015), Nunnari (2014), Richter (2014), and Anesi and Seidmann (2015), by adding experimentation and taxation components to the choice space. To establish our efficiency result for non-collegial voting rules and limited redistribution (Proposition 2), we exploit the constructive techniques developed in Baron and Bowen (2015), Anesi and Seidmann (2015), and Anesi and Duggan (2018) but push this further to construct an efficient Markovian equilibrium in a non-stationary environment where: (i) the size of the benefits allocated in each period is endogenously determined by policy choices, and (ii) the committee members' policy preferences over an additional (non-redistributive) policy dimension evolve with their endogenous beliefs. Our characterization of non-stationary equilibria for collegial voting rules also extends the work of Nunnari (2014) who focuses on stationary Markov perfect equilibria in divide-the-dollar environments with veto players.

Merlo and Wilson (1995) and Eraslan and Merlo (2002) also analyze legislative bargaining games in which the size of the revenues to be distributed among committee members evolves stochastically over time. In contrast to our framework, however, the stochastic

²A notable exception is recent work by Hirsch et al. (2015) who find that veto players can have positive effects under some conditions by forcing higher quality policy change.

³This inefficiency is also supported by some empirical work. For example, Peritz (2017) provides empirical evidence that a greater number of veto players can hinder compliance with WTO dispute rulings.

process that governs the evolution of benefits is autonomous, and bargaining ends as soon as an agreement is reached.

This paper is related to the substantial body of political economy research studying political failures, which was first articulated by Besley and Coate (1998). More closely related are papers by Alesina and Drazen (1991), Fernandez and Rodrik (1991) and Dziuda and Loeper (2016). These papers explore inefficient policy persistence, but do not consider how this might be affected by the voting rule of the committee deciding how to distribute resources. Dewatripont and Roland (1992) examined gradualism in reforms, but did not consider how this is affected by redistribution. Tornell (1998) also provides a theory of reform, but does not focus on the impact of redistribution.

The remainder of the paper is organized as follows. In Section 2, we present our model of dynamic policy making in a committee. In Sections 3 and 4, we analyze two important benchmarks – the optimal stopping rule and the equilibrium outcome in the case of no redistribution. In Section 5 we consider that redistribution is allowed and consider non-collegial and collegial rules separately. We conclude with a discussion of the results in Section 6.

2 Model

Players, policies and preferences. We present a stylized model of dynamic policy making by a committee (such as a legislature, or members of an international organization), which consists of $n \geq 3$ members: $N \equiv \{1, \dots, n\}$. Time is divided into discrete periods of length $\Delta > 0$, and the committee meets at the beginning of each period. We will subsequently focus on the limiting case as Δ becomes arbitrarily short.⁴

⁴This approach of “discretizing” dynamic games is common in the experimentation literature (e.g., Murto and Välimäki (2011)). It permits the analysis of how heterogeneous agents collectively trade off exploration and exploitation in experimentation, while avoiding the standard difficulties inherent in formulating continuous-time games with history-dependent behavior (e.g., Bergin and MacLeod (1993)). The heterogeneity of preferences over experimentation, which is crucial to our analysis, would vanish with patience in the discrete analogue of our model.

In each period t the committee has to choose a policy p^t that has three components. The first component of the policy a^t is a choice to engage in a risky reform R or implement a known safe alternative S . The reform R is either good or bad. There are two possible outcomes if it is implemented, success or failure. The probability of success is $\gamma\Delta$ if the reform is good and the probability of success is 0 if it is bad. Thus good news from the reform is conclusive. Committee member i places value $r_i > 0$ on a success, and 0 on a failure. If implemented, the safe alternative S gives individual i a per-period benefit of $\Delta s_i > 0$ with probability one. Let $\bar{r} \equiv \sum_{i \in N} r_i$ and $\bar{s} \equiv \sum_{i \in N} s_i$, and assume that $\gamma\bar{r} > \bar{s} > 0$, so that the reform, if good, is better than the safe alternative in expectation. The second component of p^t is a tax rate on individual benefits τ^t . We assume that there is an exogenous upper bound $\hat{\tau} \in [0, 1]$ on τ^t , so that $\tau^t \in [0, \hat{\tau}]$. This upper bound represents an exogenous constraint on redistribution. The third component of p^t , denoted x^t , is a choice of how to redistribute the tax revenues raised in period t and hence $x^t \in X \equiv \{(x_1, \dots, x_n) \in [0, 1]^n : \sum_{i \in N} x_i = 1\}$.

If policy $p^t = (a^t, \tau^t, x^t)$ is implemented at the start of period t and committee member i believes that the reform is good with probability α , then her (per-period) expected payoff is given by

$$w_i(a^t, \tau^t, x^t \mid \alpha) \equiv \begin{cases} \alpha\gamma\Delta[(1 - \tau^t)r_i + \tau^t x_i^t \bar{r}] & \text{if } a^t = R, \\ \Delta[(1 - \tau^t)s_i + \tau^t x_i^t \bar{s}] & \text{if } a^t = S. \end{cases}$$

Committee members discount at the continuously compounded rate ρ — so that the common discount factor is $\delta = e^{-\rho\Delta}$ — and seek to maximize their average discounted sums of payoffs.

Policy making. We model policy-making using a dynamic bargaining framework with an endogenous status quo. Each period t begins with a status quo policy p^{t-1} , inherited from the previous period. An order of proposers $(\pi_1, \dots, \pi_n)$ is randomly selected from the set Π of all permutations of N , with each permutation in Π having a positive probability of being selected. Proposer π_1 then makes the first proposal $p = (a, \tau, x) \in \{R, S\} \times [0, \hat{\tau}] \times X$; once the proposal is made committee members vote sequentially (in arbitrary order) over

whether to accept it.⁵ The proposal is accepted if a coalition $C \in \mathcal{D}$ of committee members vote to accept, and it is rejected otherwise, where $\mathcal{D} \subseteq 2^N \setminus \{\emptyset\}$ is the collection of *decisive coalitions*. If the proposal is accepted, then it is implemented, payoffs accrue and the game transitions to the next period, where the new status quo is $p^t = p$; otherwise, proposer π_2 is called upon to make a proposal and the same process is repeated. If the n proposers all make unsuccessful proposals, then the status quo p^{t-1} is implemented and remains the status quo in period $t + 1$. Each proposal round takes a negligible amount of time.

The game begins with the exogenously given status quo $p^0 \equiv (a^0, \tau^0, x^0)$, where $a^0 \equiv S$, $\tau^0 \equiv 0$ and $x_i^0 \equiv s_i/\bar{s}$ for all $i \in N$. That is, the initial status quo consists of the safe alternative and no redistribution of individual benefits.

We make standard assumptions on the set of decisive coalitions $\mathcal{D}$ (e.g., Austen-Smith and Banks (1999)). We assume the voting rule $\mathcal{D}$ is *proper*, i.e., every pair of decisive coalitions has a nonempty intersection: $C_1, C_2 \in \mathcal{D}$ implies $C_1 \cap C_2 \neq \emptyset$. Moreover, we assume $\mathcal{D}$ is *monotonic*, i.e., any superset of a decisive coalition is itself decisive: $C_1 \in \mathcal{D}$ and $C_1 \subseteq C_2$ imply $C_2 \in \mathcal{D}$. We will say that $\mathcal{D}$ is *collegial* if there is some committee member who belongs to every decisive coalition, i.e., $\bigcap \mathcal{D} \neq \emptyset$; we refer to such a committee member as a *vetoer*. The family of collegial voting rules includes unanimity rule (i.e., $\mathcal{D} = \{N\}$) and dictatorships (i.e., $\mathcal{D} = \{C \subseteq N : C \ni i\}$ for some $i \in N$). If no player has a veto, then $\mathcal{D}$ is *non-collegial*. For example, any quota rule defined by $\mathcal{D} = \{C : |C| \geq q\}$, where q satisfies $n/2 < q < n$, is non-collegial.

Learning. The initial probability that the risky reform R is good is given by $\alpha_0 \in (0, 1)$. Committee members update their (common) belief about R 's type through the sequence of policy choices using Bayes' rule. The first successful trial of R reveals to all committee members that it is good, and hence the common belief updates to one. In the event that $k \in \mathbb{N}$ trials are unsuccessful, the belief is

$$\alpha_k \equiv \frac{\alpha_0(1 - \gamma\Delta)^k}{\alpha_0(1 - \gamma\Delta)^k + 1 - \alpha_0}.$$

⁵Sequential voting is the standard approach with non-Markovian equilibrium concepts — e.g., Cho et al. (2009). This ensures that agents always vote as if pivotal.

Let $A \equiv \{\alpha_k : k = 0, 1, 2, \dots\} \cup \{1\}$ be the set of possible values for the belief.

Equilibrium. As stated, our objective is to explore the institutional mechanisms that support socially efficient experimentation. To do so, we follow closely the approach taken by Acemoglu et al. (2008), studying conditions under which efficient experimentation can be sustained by renegotiation-proof (pure strategy) perfect Bayesian equilibria (PBEs). Note that, in a PBE of this game, committee members' beliefs are necessarily given by Bayes' rule. A PBE is said to be renegotiation-proof if, after any public history, there does not exist another PBE that can make all players weakly better off (and some strictly better off).⁶ However, we will not impose any refinement of PBE when stating our negative results (those claiming that efficient policies cannot be sustained) in order to make them stronger. In order to limit the number of possible cases (without affecting the paper's conclusions), we will also assume that in case of a tie, a player will prefer to continue rather than to stop experimenting. Henceforth we will refer to a PBE with this tie-breaking rule as simply an *equilibrium*, and we will refer to a renegotiation-proof PBE with this tie-breaking rule as a *renegotiation-proof equilibrium*.

3 The Optimal Stopping Rule

To highlight the normative implications of redistribution on experimentation outcomes, we analyze the benchmark case of a utilitarian social planner to which later results can be compared. Consider the problem of a social planner whose objective is to maximize aggregate payoffs. This is a standard Markov decision problem with the planner's belief as a state variable. When the belief is α_k and the planner implements the risky reform R she obtains, in addition to the expected aggregate revenue $\alpha_k \gamma \Delta \bar{r}$, some information that she uses to update her beliefs. When she implements the safe alternative S she only obtains the aggregate revenue $\bar{s} \Delta$ and her belief remains unchanged. Therefore, if the optimal solution requires S to be implemented in a given period t , then the belief will remain the

⁶Though the game has imperfect information, all players are symmetrically informed at every history and, therefore, the usual interpretation of renegotiation-proofness applies.

same and S will also be implemented in all future periods. As is standard in the literature on experimentation, it follows that the optimal solution is a stopping rule: there exists a $k^* \in \mathbb{N}$ such that, after k^* unsuccessful trials of R , the belief is so low that the planner suspends experimentation and implements the safe alternative only.

Formally, let $V^*(\alpha)$ be the planner's average discounted value from a period that begins with a belief α . That is, if the planner applies the optimal stopping rule, then $V^*(\alpha)$ is the expected sum of the committee members' average discounted payoffs from the resulting outcome path. As $\gamma\bar{r} > \bar{s}$, we evidently have $V^*(1) = \gamma\Delta\bar{r}$. If the social planner chooses to continue experimenting when the belief is α_k , then her expected payoff is equal to $[1 - \delta(1 - \gamma\Delta)]\alpha_k\gamma\Delta\bar{r} + \delta(1 - \alpha_k\gamma\Delta)V^*(\alpha_{k+1})$. If she chooses to stop, then her payoff is $\bar{s}\Delta$. Hence,

$$V^*(\alpha_k) = \max \left\{ [1 - \delta(1 - \gamma\Delta)]\alpha_k\gamma\Delta\bar{r} + \delta(1 - \alpha_k\gamma\Delta)V^*(\alpha_{k+1}), \bar{s}\Delta \right\}.$$

Denote the optimal cutoff by α^* . Recall that periods are discrete and thus α^* must be an element of the set of feasible beliefs A . The optimal cutoff α^* may thus be strictly smaller than the belief that makes the social planner indifferent between continuing with the reform for one more period and switching to the safe alternative. Noting that if $\alpha_k = \alpha^*$, then $V^*(\alpha_{k+1}) = \bar{s}\Delta$, we obtain

$$\alpha^* \equiv \min \left\{ \alpha_0, \max \left\{ \alpha \in A : \alpha < \frac{(1 - \delta)}{\gamma[(1 - \delta(1 - \gamma\Delta))(\bar{r}/\bar{s}) - \delta\Delta]} \right\} \right\}.$$

We will say that an equilibrium *sustains the optimal stopping rule* if the committee applies the optimal stopping rule on the equilibrium path.

To make things interesting, we wish to study situations where social optimality dictates to experiment for at least one period, and thus $\alpha^* < \alpha_0$, for Δ sufficiently small. Note that, as $\Delta \rightarrow 0$, the social planner's ideal cutoff converges to

$$\min \left\{ \alpha_0, \frac{\rho}{\gamma[(\rho + \gamma)(\bar{r}/\bar{s}) - 1]} \right\}.$$

Imposing $\alpha_0 > \rho/(\gamma[(\rho + \gamma)(\bar{r}/\bar{s}) - 1])$ guarantees that there is a $\hat{\Delta} > 0$ such that, for all $\Delta < \hat{\Delta}$, we have $\alpha^* < \alpha_0$. And thus some experimentation is optimal for Δ sufficiently small. Throughout, we maintain this assumption.

Assumption A1. $\alpha_0 > \rho/(\gamma[(\rho + \gamma)(\bar{r}/\bar{s}) - 1])$.

4 Policy Experimentation without Redistribution

We now return to the analysis of the bargaining game introduced in Section 2. Throughout this section, we assume that no redistribution is permitted, i.e., $\hat{\tau} = 0$. Comparison of this case to the social planner's benchmark reveals that in the absence of redistribution, socially efficient experimentation typically cannot be sustained in equilibrium.

To begin we must establish some notation. By the same logic as above, each committee member i 's ideal experimentation plan is a stopping rule with cutoff

$$\hat{\alpha}_i \equiv \min \left\{ \alpha_0, \max \left\{ \alpha \in A : \alpha < \frac{(1 - \delta)}{\gamma[(1 - \delta(1 - \gamma\Delta))(r_i/s_i) - \delta\Delta]} \right\} \right\} .$$

This is the rule that committee member i would implement in every equilibrium if she were a dictator, i.e., if $\mathcal{D} = \{C \subseteq N : C \ni i\}$. Note that $\hat{\alpha}_i$ is (weakly) decreasing in the ratio r_i/s_i . That is, each committee member's incentive to experiment increases with the extent to which she values the reform over the safe alternative. Henceforth, without loss of generality, we order committee members such that $r_i/s_i \leq r_{i+1}/s_{i+1}$ and thus committee member 1 wishes to cease experimenting first. Moreover, we refer to $\hat{\alpha}_i$ as committee member i 's *ideal cutoff*, and to r_i/s_i as her *benefit ratio*. We find it useful to identify two key members of the committee defined as follows.⁷

Definition 1. *Committee member $i \in N$ is a pivot for voting rule $\mathcal{D}$ if $\{j \in N : \hat{\alpha}_j < \hat{\alpha}_i\} \notin \mathcal{D}$ and $\{j \in N : \hat{\alpha}_j > \hat{\alpha}_i\} \notin \mathcal{D}$. The set of pivots for $\mathcal{D}$ is denoted $P(\mathcal{D})$, and we refer to $\mathbb{l} \equiv \min P(\mathcal{D})$ and $\mathbb{r} \equiv \max P(\mathcal{D})$ as the left and right pivots, respectively.*

Observe that coalitions $\{1, \dots, \mathbb{r}\}$ and $\{\mathbb{l}, \dots, n\}$ must be decisive for every voting rule $\mathcal{D}$; otherwise, by monotonicity of $\mathcal{D}$, the right pivot would be greater than $\mathbb{r}$ and the left pivot smaller than $\mathbb{l}$. Furthermore, it follows immediately from Definition 1 that coalition

⁷The terms *left pivot* and *right pivot* in Definition 1 are related to similar terms described in Dziuda and Loeper (2018). We provide a formal definition for our context.

$\{1, \dots, \mathbb{I}\}$ is *blocking*, in the sense that a policy proposal can only be accepted if at least one member of that coalition votes to accept it.

Next, let $\bar{V}_i(\alpha_k)$ be the average dynamic payoff to committee member i at the start of the game induced by a stopping rule with cutoff $\alpha_k \in A \setminus \{1\}$. When the cutoff is α_k , in each of the first k periods, player i receives $(1 - \delta)r_i$ if the reform is good and successful. This occurs with probability $\alpha_0\gamma\Delta$. Starting from period $k + 1$ onward there are two possible cases. In the first case, the reform succeeds in at least one of the first k periods, and it is implemented from period $k + 1$ on. This occurs with probability $\alpha_0[1 - (1 - \gamma\Delta)^k]$ and yields a per-period expected payoff of $(1 - \delta)\gamma\Delta r_i$. In the second case, the first k trials are all unsuccessful, and the safe alternative S is implemented from period $k + 1$ on. This occurs with the complementary probability $[1 - \alpha_0(1 - (1 - \gamma\Delta)^k)]$ and yields a per-period payoff of $(1 - \delta)\Delta s_i$. Thus committee member i 's dynamic payoff for cutoff α_k is given by

$$\bar{V}_i(\alpha_k) \equiv (1 - \delta^k)\alpha_0\gamma\Delta r_i + \delta^k \left[\alpha_0[1 - (1 - \gamma\Delta)^k]\gamma\Delta r_i + [1 - \alpha_0(1 - (1 - \gamma\Delta)^k)]\Delta s_i \right],$$

which simplifies to

$$\bar{V}_i(\alpha_k) = [1 - \delta^k(1 - \gamma\Delta)^k]\alpha_0\gamma\Delta r_i + \delta^k[1 - \alpha_0(1 - (1 - \gamma\Delta)^k)]\Delta s_i. \quad (1)$$

Our first result gives a complete characterization of the equilibria for the bargaining game in terms of the voting rule and the distribution of ideal cutoffs.

Proposition 1. *Suppose $\hat{\tau} = 0$. For all voting rules there is a unique equilibrium outcome that takes the form of a stopping rule with cutoff:*

$$\bar{\alpha} = \begin{cases} \hat{\alpha}_r & \text{if } \bar{V}_l(\hat{\alpha}_r) \geq s_l\Delta, \\ \alpha_0 & \text{otherwise.} \end{cases}$$

In words, Proposition 1 states that if experimentation begins it always ends at the right pivot's ideal cutoff $\hat{\alpha}_r$. Furthermore, experimentation only begins if the left pivot's dynamic payoff using this cutoff, $\bar{V}_l(\hat{\alpha}_r)$, exceeds the left pivot's status quo payoff, $s_l\Delta$. Otherwise, experimentation will never begin. Thus the left pivot determines whether experimentation begins, and the right pivot determines the belief at which experimentation ends.

An immediate consequence of Proposition 1 is that if n is odd and the voting rule is simple majority, i.e. $\mathcal{D} = \{C \subseteq N: |C| \geq (n+1)/2\}$, then the stopping rule with median committee member's ideal cutoff is implemented. This is reminiscent of the well-known median voter theory, where an odd number of voters with static, single-peaked preferences must collectively choose a policy from a unidimensional choice space. However, the logic behind Proposition 1 is more subtle, not only because this is a dynamic setting with evolving status quo and beliefs, but mainly because policy preferences in each period are *endogenously* determined by equilibrium behavior in future periods. To see this, we discuss the intuition for the proof below.

Although equilibrium continuation values could intricately depend on the previous history of play, we first show that if the belief becomes smaller than or equal to committee member n 's ideal cutoff, $\hat{\alpha}_n$, the safe alternative must be implemented in every period of every continuation game. Therefore, when the belief is α_k with $\alpha_{k+1} = \hat{\alpha}_n$, the members of the decisive coalition $\{1, \dots, r\}$ can play in accordance with their own preferences without risking to trigger adverse decisions in future periods: they will always agree to switch from R to S if R is the status quo, and will always reject any proposal to change S to R if S is the status quo.

Applying the same logic recursively, we obtain that whenever the belief α_k is smaller than $\hat{\alpha}_r$, the equilibrium outcome of every continuation game is unique and has to be a stopping rule with cutoff α_k . Now consider any belief $\alpha_k \geq \hat{\alpha}_r$. All agents that wish to experiment more prefer to do so right away rather than wait to experiment because by waiting, expected future benefits are further discounted. This turns the dynamic bargaining problem into the choice between two options: implementing the stopping rule with cutoff $\hat{\alpha}_r$, or maintaining the safe alternative. If $\bar{V}_l(\hat{\alpha}_r) \geq s_l \Delta$, then the left pivot l and, by single-peakedness, all the other members of the decisive coalition $\{l, \dots, n\}$ prefer the first option. The stopping rule with cutoff $\hat{\alpha}_r$ must therefore be the unique equilibrium outcome. If instead $\bar{V}_l(\hat{\alpha}_r) < s_l \Delta$, then the left pivot l and all the other members of the blocking coalition $\{1, \dots, l\}$ prefer to maintain the initial status quo S and, consequently, experimentation never occurs — i.e., the unique equilibrium cutoff is α_0 .

It follows from the previous discussion that with simple majority voting and an odd number of heterogeneous committee members, the committee always over-experiments if $\hat{\alpha}_r < \alpha^*$, and always under-experiments if $\hat{\alpha}_r > \alpha^*$. With any voting rule, it over-experiments if $\hat{\alpha}_r < \alpha^*$ and $\bar{V}_l(\hat{\alpha}_r) \geq s_l \Delta$, and it under-experiments either if $\hat{\alpha}_r > \alpha^*$ and $\bar{V}_l(\hat{\alpha}_r) \geq s_l \Delta$, or if $\bar{V}_l(\hat{\alpha}_r) < s_l \Delta$ and $\alpha^* \neq \alpha_0$. Assumption A1 guarantees that $\alpha^* > \alpha_0$ for Δ sufficiently small and hence some experimentation is optimal. We summarize this in the following corollary.

Corollary 1. *Suppose $\hat{\tau} = 0$. If $r_r/s_r \neq \bar{r}/\bar{s}$, then there exists $\Delta_0 > 0$ such that, whenever $\Delta < \Delta_0$, all equilibria fail to sustain the optimal stopping rule.*

Corollary 1 states that, in the absence of redistribution, every equilibrium fails to sustain the optimal stopping rule (in the limit as $\Delta \rightarrow 0$) whenever the right pivot's benefit ratio differs from $\bar{r}/\bar{s}$, which is generically the case. We conclude that efficient experimentation is typically impossible without redistribution.

We end this section with a remark on Pareto inefficiency. Proposition 1 and Corollary 1 show how under-experimentation may happen under any voting rule, yielding socially inefficient outcomes in equilibrium (in a Utilitarian sense). But such equilibrium outcomes may even be Pareto dominated. This is illustrated by the following example.

Suppose $n = 5$ and $\mathcal{D}$ is a quota rule with quota $q = 4$ (so that $\mathbb{I} = 2 < 4 = r$). Set $\rho = \gamma = \alpha_0 = 2/3$, $r_i = 2$ for all i , $s_i = 1$ for all $i \in \{1, 2, 3\}$, and $s_i = \epsilon$ for all $i \in \{4, 5\}$, where $\epsilon > 0$ is arbitrarily small. It is readily checked that under these assumptions, $\lim_{\Delta \rightarrow 0} \hat{\alpha}_i = 3/5 < \alpha_0$ for all $i \in \{1, 2, 3\}$, and $\lim_{\Delta \rightarrow 0} \hat{\alpha}_i = 3\epsilon/(8 - 3\epsilon)$ for all $i \in \{4, 5\}$. Hence, for arbitrarily small $\Delta > 0$, the right pivot's optimal stopping rule converges to perpetual experimentation as $\epsilon \rightarrow 0$. Further, $\bar{V}_l(\hat{\alpha}_r) \rightarrow 8\Delta/9 < \Delta = s_l \Delta$, as $\epsilon \rightarrow 0$, so the left pivot's dynamic payoff under the right pivot's optimal stopping rule is lower than the left pivot's status quo payoff. It therefore follows from Proposition 1 that, for sufficiently small ϵ (and sufficiently small $\Delta > 0$), the reform is never implemented in equilibrium, although all committee members would be better off experimenting for a positive number of periods. The reason is that, while committee members in the blocking coalition $\{1, 2\}$ would like to experiment, because of the endogeneity of the status quo,

they fear that experimentation will go on for too long, so they prefer not to experiment at all.

5 Policy Experimentation and Redistribution

We have seen in Section 4 that the optimal stopping rule is typically not sustainable in equilibrium without redistribution. In this section, we ask whether this is still true when it is possible to redistribute the revenues from experimentation among committee members. We begin with the benchmark case in which the committee can freely redistribute revenues among its members, and then turn to the (more empirically relevant) case of constrained redistribution.

5.1 Policy Experimentation with Unconstrained Redistribution

Suppose $\hat{\tau} = 1$, so that redistribution is unconstrained. As all resources can be redistributed, benefit ratios are no longer relevant and each agent that receives a share of the surplus has an incentive to maximize that share. This is maximized when the optimal stopping rule is implemented. Moreover, full redistribution allows great flexibility in creating “rewards” and “punishments” that provide the incentives to implement optimal policies. Yet, we must ensure that any such punishments are themselves consistent with renegotiation-proofness. Despite this, we can show that every renegotiation-proof equilibrium sustains the optimal stopping rule for sufficiently small period length.

Proposition 2. *Suppose $\hat{\tau} = 1$. Then, for any voting rule, there exists $\bar{\Delta} > 0$ such that, for all $\Delta < \bar{\Delta}$, renegotiation-proof equilibria exist and all of them sustain the optimal stopping rule.*

The proof of this result borrows insights from the repeated-games literature. For every voting rule $\mathcal{D}$, we first identify a lower-bound $\underline{w}_i(p \mid \alpha)$ on each committee member i ’s equilibrium payoff at the start of any period that begins with status quo policy p and belief α . It follows that equilibrium payoff vectors must belong to the simplex $W(p \mid \alpha) \equiv \{(w_1, \dots, w_n) \in \mathbb{R}^n : \sum_{i=1}^n w_i \leq V^*(\alpha), \text{ and } w_i \geq \underline{w}_i(p \mid \alpha) \text{ for all } i\}$. Recall that p^0 is

the status quo policy in period zero with the safe alternative and no redistribution. Next, for each payoff vector w^* on the Pareto frontier of $W(p^0 \mid \alpha_0)$ we construct a renegotiation-proof equilibrium that supports w^* (for sufficiently small Δ), thus establishing that every renegotiation-proof equilibrium payoff vector $(w_1, \dots, w_n)$ must satisfy $\sum_{i=1}^n w_i = V^*(\alpha_0)$, as desired. In this equilibrium construction, committee members are able to coordinate on an efficient experimentation rule not only at the beginning of the game but also at any time during the game, whether on or off the equilibrium path. Should any member i ever deviate from the prescribed behavior at any status quo p and belief α , the committee implements an efficient equilibrium point in $W(p \mid \alpha)$ that gives i her most severe punishment payoff $\underline{w}_i(p \mid \alpha)$.

Thus, without bounds on redistribution, collective experimentation must yield efficient outcomes. In view of Corollary 1, this result may give the impression that constraints on redistribution are the primary cause of inefficiency in collective experimentation. In the next section we will see that this is not the complete picture for the realistic case of constrained redistribution. There we see that voting rules may play a more significant role.

5.2 Policy Experimentation with Constrained Redistribution

Proposition 2 prompts the following question: How much stronger is full redistribution than the condition necessary for efficient experimentation in equilibrium? The answer critically turns on the voting rule. We first consider non-collegial voting rules and then collegial rules.

5.2.1 Non-collegial Voting Rules

In stark contrast to the case of no redistribution, our next proposition states that with non-collegial voting rules the optimal stopping rule can be sustained by a renegotiation-proof equilibrium for *any* level of redistribution. In fact, something stronger is true: although stationary Markov strategies sharply constrain the ability to punish and reward committee members for past behavior, the renegotiation-proof equilibrium that sustains the optimal

stopping rule can be taken to be stationary Markov.⁸

Proposition 3. *Suppose $\mathcal{D}$ is non-collegial, i.e., no committee member has a veto. Then, the following holds for every upper bound $\hat{\tau} > 0$: there exists $\bar{\Delta} > 0$ such that, for all $\Delta < \bar{\Delta}$, the optimal stopping rule is sustained by a (stationary Markov) renegotiation-proof equilibrium.*

The proof of Proposition 3 is constructive. The idea of the construction is simple: it is built around a collection of $(n - 1)$ -member coalitions — the potential “*governing coalitions*.” These governing coalitions have the property that each committee member belongs to at least one of these coalitions but not to all of them. The same thing inevitably happens from any history (both on and off the equilibrium path): the optimal stopping rule is implemented and, in every period, the members of a given governing coalition equally share the sum of the expected aggregate revenues that can be redistributed, i.e., $\hat{\tau}\alpha_k\gamma\Delta\bar{r} > 0$ if the belief α_k exceeds α^* , and $\hat{\tau}\bar{s}\Delta > 0$ otherwise.

At the start of every period, the status quo policy and the current belief — which are payoff relevant — reveal to the committee whether play in the previous period was consistent with the optimal stopping rule, and whether a governing coalition formed (i.e., equally shared the entire transferable benefits among its members). If this is the case, then the same governing coalition forms again and continues to implement the optimal stopping rule; otherwise, the (randomly selected) first proposer successfully offers to form a governing coalition and to follow the optimal stopping rule.

Given the inevitability of this process, the best possible scenario for any committee member is to form or be a member of the governing coalition that will share the transferable revenues from experimentation in every future period. For any member i of such a coalition, the potential benefits of a one-period deviation vanish as the period length Δ becomes arbitrarily small, whereas the long-run cost does not: a deviation would trigger the formation of a new governing coalition in the next period, which i might not be a

⁸Of course, renegotiation-proofness already constitutes an obstacle to the construction of efficient equilibria, as it may reduce the severity of the off-path “punishments” available to support the appropriate incentives (which must themselves be renegotiation-proof).

member of. Though this would have no impact on the proportion $(1 - \hat{\tau})$ of her future payoffs that cannot be redistributed (since the optimal stopping rule is implemented in any case), she would potentially lose a proportion $\hat{\tau} > 0$ of her share of future aggregate revenues as a member of the governing coalition. As no member of a governing coalition is prepared to run such a risk and governing coalitions are decisive (recall that no player has a veto), profitable deviations from the prescribed path are impossible. Moreover, this equilibrium is renegotiation-proof since it generates payoff vectors on the Pareto frontier both on and off the path.

In contrast to Corollary 1, which is an “impossibility result”, Proposition 3 says that it is possible for the committee to implement the optimal stopping rule in equilibrium. In line with the relational-contracts literature, a possible interpretation of the proposition is that committee members can credibly agree to play an equilibrium that sustains efficient experimentation, and could not agree to replace it with another equilibrium once it is in place. While we follow Acemoglu et al.’s (2008) approach of focusing on the “best sustainable mechanisms,” it should be noted, however, that there are other renegotiation-proof equilibria that do not sustain the optimal stopping rule. Therefore, it is not clear how committee members would select an equilibrium. One can argue that the above equilibrium is focal in the sense of Schelling (1960), as it is the only one whose payoff vector belongs to the (unconstrained) Pareto frontier; this is the approach taken, for example, by Dixit and Olson (2000).

5.2.2 Collegial Voting Rules

Given the result obtained for non-collegial voting rules in the previous subsection, it is natural to ask whether the optimal stopping rule is also sustainable with any level of redistribution under collegial voting rules. A little reflection suggests that the answer is no because of the presence of veto players. Indeed, efficient experimentation requires two sets of incentive constraints to be met. The first set ensures that the “vetoers” all agree to change the initial status quo policy to some (R, τ, x) , $x \in X$ and $\tau \in [0, \hat{\tau}]$; the second ensures that if the belief becomes equal to α^* , then they all agree to stop experimenting.

Formally, the first constraint requires that each vetoer i 's equilibrium continuation value from implementing (R, τ, x) in the first period is greater than or equal to her payoff from maintaining the initial status quo in all future periods. In the benchmark case where $\hat{\tau} = 0$, this is equivalent to $\bar{V}_i(\alpha^*) \geq \Delta s_i$ for each committee member i who has a veto (where $\bar{V}_i(\cdot)$ is given in equation (1)). We rearrange terms, take the limit as Δ goes to zero, and apply l'Hôpital's rule to obtain

$$[1 - e^{\psi(1+\rho\gamma)}] \alpha_0 \gamma r_i + e^{\psi\rho\gamma} (1 - \alpha_0 + e^{\psi} \alpha_0) s_i \geq s_i ,$$

where $\psi \equiv \log \left[\frac{\rho \bar{s}}{(\rho + \gamma)(\gamma \bar{r} - \bar{s})} \right] < 0$.

Let $\underline{v}$ be the vetoer who wishes to cease experimenting first, i.e., $\underline{v} \equiv \min \bigcap \mathcal{D}$. It follows that if $\underline{v}$'s benefit ratio satisfies

$$\frac{r_{\underline{v}}}{s_{\underline{v}}} < \frac{1 - e^{\psi\rho\gamma} (1 - \alpha_0 + e^{\psi} \alpha_0)}{[1 - e^{\psi(1+\rho\gamma)}] \alpha_0 \gamma} , \quad (2)$$

then there exists $T > 0$ such that, for every $\hat{\tau} \in [0, T)$, her incentive constraint is always violated for arbitrarily small Δ 's. Intuitively, when $\hat{\tau} < T$, the permitted level of redistribution is not sufficiently large to compensate committee member $\underline{v}$'s loss from experimenting and, consequently, the optimal stopping rule is not sustainable in equilibrium.

The second key incentive constraint in the construction of efficient equilibria formally requires that if the belief becomes equal to α^* , then each vetoer i 's continuation value from accepting a proposal to cease experimenting exceeds her equilibrium continuation value from rejecting it. It is readily checked that the former continuation value can be written as $[(1 - \hat{\tau})s_i + \hat{\tau}y_i\bar{s}]\Delta$, for some $y \in X$.⁹ For every status quo policy (R, τ', z) , the latter continuation value is bounded below by $\alpha^* \gamma \Delta [(1 - \tau')r_i + \tau'z_i\bar{r}]$. Hence, an obvious necessary condition for the existence of an equilibrium that supports the optimal stopping rule is that $(1 - \hat{\tau})s_i + \hat{\tau}y_i\bar{s} \geq \alpha^* \gamma [(1 - \tau')r_i + \tau'z_i\bar{r}]$ for some $y, z \in X$, $\tau' \in [0, \hat{\tau}]$, and all

⁹If policy (S, τ^t, z^t) is implemented in period t , then i receives a payoff of $[(1 - \hat{\tau})s_i + \hat{\tau}y_i^t\bar{s}]\Delta$, where $y_i^t \equiv [1 - (\tau^t/\hat{\tau})](s_i/\bar{s}) + (\tau^t/\hat{\tau})z_i^t$. Letting y be the average discounted distribution of tax revenues from the first period the committee stopped experimenting on, we obtain a continuation value of $[(1 - \hat{\tau})s_i + \hat{\tau}y_i\bar{s}]\Delta$ for each i .

vetoers i . Setting $\hat{\tau} = 0$ and letting Δ go to zero, we can rewrite this condition as

$$s_i \geq \frac{\rho r_i \bar{s}}{(\rho + \gamma) \bar{r} - \bar{s}}.$$

Let $\bar{v}$ be the vetoer who wishes to cease experimenting last, i.e., $\bar{v} \equiv \max \bigcap \mathcal{D}$. If the latter's benefit ratio is large relative to the social planner's, i.e., if

$$(\rho + \gamma) \frac{\bar{r}}{\bar{s}} < 1 + \rho \frac{r_{\bar{v}}}{s_{\bar{v}}}, \quad (3)$$

then there exists $T > 0$ such that, for every $\hat{\tau} \in [0, T)$ and arbitrarily small Δ , vetoer $\bar{v}$'s incentive constraint cannot hold. Because of the committee's imperfect ability to redistribute the benefits from ending experimentation towards vetoer $\bar{v}$, the latter must reject any proposal to stop experimenting when the belief is α^* in equilibrium.

The discussion above is summarized in the following proposition.

Proposition 4. *Suppose $\mathcal{D}$ is collegial (i.e., some players have a veto). If either condition (2) or (3) is satisfied, then there exists $T > 0$ such that the following holds for all $\hat{\tau} \in [0, T)$: there is $\hat{\Delta} > 0$ such that, whenever $\Delta < \hat{\Delta}$, every equilibrium fails to sustain the optimal stopping rule.*

This proposition shows that institutional details matter: in contrast to any non-collegial voting rule, *efficient experimentation may not be attainable under a collegial voting rule if not enough redistribution is permitted*. Under the premises of the proposition, even unrefined equilibria all fail to support the optimal stopping rule. Moreover, there is no general lower bound on $\hat{\tau} < 1$ allowing to avoid this negative conclusion. This is easily seen by considering the case where $\mathcal{D}$ is dictatorial, i.e., $\mathcal{D} = \{C \subseteq N : C \ni i\}$ for some $i \in N$. It follows immediately from the analysis in Section 3 that if the dictator's benefit ratio differs from the social planner's (i.e., if $r_i/s_i \neq \bar{r}/\bar{s}$) and $\hat{\tau} < 1$, then every equilibrium fails to sustain the optimal stopping rule.

Coupled with Proposition 3, Proposition 4 suggests that veto rights, instead of constraints on redistribution, can hamper optimal experimentation in committees: constraints on redistribution (as long as they do not completely prevent redistribution) can only create inefficiencies if the voting rule is collegial.

6 Concluding Remarks

We analyze a dynamic model of committee decision making in which the benefits of reform may or may not be redistributed. Several conclusions emerge from the analysis. First, except for nongeneric cases, socially efficient experimentation necessitates redistribution. That is, when no redistribution is possible, the optimal stopping rule is generically unachievable. Second, if the committee can freely redistribute all the revenues from experimentation among its members, then the inefficiencies of the no-redistribution case vanish: it always implements the optimal stopping rule. Third, even in the more realistic case where redistribution is constrained, arbitrarily small amounts of redistribution suffice to support socially efficient experimentation if the voting rule is collegial. This applies to committees such as legislatures governed by majority or supermajority rules. And fourth, if the voting rule is non-collegial, then the optimal stopping rule can be sustained only with a sufficient amount of redistribution. Collegial voting rules, such as unanimity, are ubiquitous in international policymaking and other settings.¹⁰

Two implications directly follow from these results. The first speaks to a major question in the study of political institutions (e.g., McCarty (2000)): How do veto powers affect policy outcomes? It follows from our analysis that, in situations where committees can only partially redistribute the gains from policy experimentation among their members, veto powers may be detrimental to social welfare by rendering efficient experimentation unachievable. Our results also speak to the normative question of what voting rules can support efficient collective experimentation (e.g. Strulovici (2010)). Proposition 3 provides a simple answer: As long as some share of aggregate revenues can be redistributed, efficient experimentation can be sustained under *any* non-collegial voting rule.

The redistributive tool we have used to support efficient policy experimentation can be used in other settings. In the case of non-collegial voting rules we have supported efficient experimentation by linking the equilibrium of the experimentation game to an equilibrium of the redistribution game with a structure of dynamic coalitions (similar to Anesi and

¹⁰Maggi and Morelli (2006) provide theoretical conditions under which this is true if agents can choose the voting rule ex-ante.

Seidmann (2015) and Baron and Bowen (2015)). Such “linking” of policies can be used in other settings to support efficient policy making, and, indeed, has been used in practice with, for example, the Paris Climate Agreement and the Green Climate Fund.

Our analysis focuses on the frictions created by constraints on redistribution in collective experimentation. Following the previous literature, we have assumed that the committee is able to reconsider previous agreements arbitrarily frequently. It would be interesting for future research to investigate how bargaining frictions (captured with values of Δ bounded away from zero) affect collective-experimentation outcomes. This would require future work to overcome technical difficulties that have not yet been addressed in the literature on dynamic bargaining with an endogenous status quo (even without experimentation).

Appendix

A Proof of Proposition 1

We assume throughout this section that $\hat{\tau} = 0$, so that distributions have no impact on players' payoffs. To lighten the notation, we simply represent policies as elements of $\{R, S\}$, omitting the irrelevant sharing-rule component. To prove Proposition 1, it is useful to consider a class of games $\{\Gamma(a, \alpha) : a \in \{R, S\} \text{ \& } \alpha \in A\}$, where $\Gamma(a, \alpha)$ is the same game as that described in Section 2, except that it begins with an initial status quo alternative a (possibly equal to R) and a probability α (possibly different from α_0) that Alternative R is good.

It is useful to define the set of losers from the reform as $L \equiv \{i \in N : \gamma r_i < s_i\}$ and the set of winners as $W \equiv \{i \in N : \gamma r_i \geq s_i\}$.

Lemma A1. *Suppose $\hat{\tau} = 0$. Then,*

- (i) $\Gamma(R, 1)$ has a unique equilibrium outcome: Alternative R is implemented in every period if $L \notin \mathcal{D}$, and alternative S is implemented in every period otherwise;*
- (ii) for all $a \in \{R, S\}$ and $\alpha_k \leq \hat{\alpha}_n$, $\Gamma(a, \alpha_k)$ has a unique equilibrium outcome: Alternative S is implemented in every period.*

The proof of Lemma A1 can be found in the supplementary appendix. We now return to the proof of the main proposition. Suppose first that $L \notin \mathcal{D}$. Having characterized the unique PBE outcome of $\Gamma(a, 1)$ and $\Gamma(a, \alpha_k)$, for all $a \in \{R, S\}$ and all $\alpha_k \leq \hat{\alpha}_n$, we begin with an inductive argument. Take any $k \in \mathbb{N}$ such that (i) for each $a \in \{R, S\}$, alternative S is implemented in every period in any PBE of $\Gamma(a, \alpha_{k+1})$, and (ii) $\alpha_k \leq \hat{\alpha}_r$. (From Lemma A1(ii), we already know that this is the case if $\alpha_k \leq \hat{\alpha}_n$.) In the first period of $\Gamma(S, \alpha_k)$, the expected payoff to committee member i in any PBE must be a convex combination of $f_i(\alpha_k) \equiv [1 - \delta(1 - \gamma\Delta)]\alpha_k\gamma\Delta r_i + \delta(1 - \alpha_k\gamma\Delta)s_i\Delta$ and $s_i\Delta$. It follows from the definition of $\hat{\alpha}_i$, $s_i\Delta > f_i(\alpha_k)$ if and only if $\alpha_k \leq \hat{\alpha}_i$. By the same logic as in the proof of Lemma A1(i), this implies that any proposal to change status quo S to the risky alternative R must be rejected by the members of the decisive coalition $\{1, \dots, r\}$.

(The latter coalition must be decisive; otherwise, monotonicity of $\mathcal{D}$ would imply that the right pivot is greater than r .) Hence, $s_i\Delta$ is committee member i 's unique equilibrium payoff in $\Gamma(S, \alpha_k)$. It follows that committee member i 's payoff in the continuation game $\Gamma(R, \alpha_k)$ is $f_i(\alpha_k)$ if R is implemented in the first period, and $s_i\Delta$ otherwise. This implies that every member of the decisive coalition $\{1, \dots, r\}$ must accept any proposal to amend R to S (when decisive) and, therefore, at least one proposer must successfully propose S in the first period in equilibrium. We have thus established that for all $\alpha_k \leq \hat{\alpha}_r$, S is implemented in every period of $\Gamma(R, \alpha_k)$ in any PBE. The unique equilibrium outcome of $\Gamma(R, \hat{\alpha}_r)$ is therefore the stopping rule with cutoff $\hat{\alpha}_r$.

In Section 3, we defined $\bar{V}_i(\alpha_k)$ as the expected payoff to committee member i induced by the stopping rule with cutoff α_k in $\Gamma(S, \alpha_0)$. For every $0 \leq \ell \leq k$, we can similarly define the expected payoff to committee member i induced by this stopping rule in $\Gamma(S, \alpha_\ell)$ as

$$\bar{V}_i(\alpha_k | \alpha_\ell) \equiv [1 - \delta^{k-\ell}(1 - \gamma\Delta)^{k-\ell}] \alpha_\ell \gamma \Delta r_i + \delta^{k-\ell} [1 - \alpha_\ell + (1 - \gamma\Delta)^{k-\ell} \alpha_\ell] s_i \Delta.$$

Differentiating the right side of the above equation with respect to k reveals that $\bar{V}_i(\alpha_k | \alpha_\ell)$ is single-peaked in k and, therefore, in the cutoff α_k : it decreases with α_k if $\alpha_k < \hat{\alpha}_i$, and increases with α_k if $\alpha_k > \hat{\alpha}_i$. Now take any belief $\alpha_\ell > \hat{\alpha}_r$ such that the unique equilibrium outcome of $\Gamma(R, \alpha_{\ell+1})$ is the stopping rule with cutoff $\hat{\alpha}_r$. (From the previous paragraph, this is the case if $\alpha_{\ell+1} = \hat{\alpha}_r$.) The expected payoff to committee member i in any PBE of $\Gamma(R, \alpha_\ell)$ is a convex combination between $\bar{V}_i(\hat{\alpha}_r | \alpha_\ell)$ and $s_i\Delta$. Moreover, it follows from the single-peakedness of $\bar{V}_i(\cdot | \alpha_\ell)$ that $\bar{V}_i(\hat{\alpha}_r | \alpha_\ell) > s_i\Delta = \bar{V}_i(\alpha_\ell | \alpha_\ell)$, for all $i \geq r$ — recall that $\hat{\alpha}_i \leq \hat{\alpha}_r < \alpha_\ell$ for all $i \geq r$. Every member of the blocking coalition $\{r, \dots, n\}$ therefore rejects any proposal to amend the status quo R to S in the first period of $\Gamma(R, \alpha_\ell)$. This shows in particular that the unique PBE outcome of $\Gamma(R, \alpha_1)$ is the stopping rule with cutoff $\hat{\alpha}_r$.

Finally, consider the first period of the game — i.e., the first period of $\Gamma(S, \alpha_0)$. It follows from the previous paragraph that, in any equilibrium, committee member i 's payoff must be a convex combination between $\bar{V}_i(\hat{\alpha}_r) = \bar{V}_i(\hat{\alpha}_r | \alpha_0)$ and $s_i\Delta = \bar{V}_i(\alpha_0 | \alpha_0)$. Suppose first that $\bar{V}_1(\hat{\alpha}_r) < s_1\Delta$, so that $\bar{V}_i(\hat{\alpha}_r) < s_i\Delta$ for every member i of the blocking

coalition $\{1, \dots, \mathbb{l}\}$. Every member of this coalition must therefore reject any proposal to amend the status quo S to R in the first period. This in turn implies that the stopping rule with cutoff α_0 is the unique PBE outcome. Suppose now that $\bar{V}_1(\hat{\alpha}_r) \geq s_1 \Delta$. Denoting the supremum of the PBE payoffs of each committee member i by $U_i^{\sup}$, we thus have $\bar{V}_i(\hat{\alpha}_r) \geq U_i^{\sup} \geq s_i \Delta$ for every member i of the winning coalition $C \equiv \{1, \dots, n\}$. This implies that, in any PBE, the payoff of each committee member $i \in C$ from accepting a proposal to amend status quo S to R (when decisive) in the first period of the game — i.e. $\bar{V}_i(\hat{\alpha}_r)$ — must therefore exceed her payoff from rejecting it — $(1 - \delta)\Delta s_i + \delta U_i^{\sup}$. Hence, at least one member of the committee must successfully propose alternative R in the first period. This in turn implies that the stopping rule with cutoff $\hat{\alpha}_r$ is the unique PBE outcome, completing the proof of Proposition 1.

B Proof of Proposition 2

In this section we establish prove Proposition 2 for the case of collegial rules. The proofs for the various collegial rules are relegated to the supplementary appendix.

Suppose that $\bigcap \mathcal{D} = \emptyset$. Let $\bar{\Delta} \equiv \sup \{ \Delta \in \mathbb{R}_+ : (1 - e^{-\rho\Delta})\gamma\bar{r} < e^{-\rho\Delta}\bar{s}/(n - 1) \} > 0$. Observe for future reference that if $\Delta < \bar{\Delta}$, then

$$(1 - \delta)|w_i(p \mid \alpha) - w_i(p' \mid \alpha)| \leq (1 - \delta)\gamma\bar{r}\Delta < \delta \frac{\bar{s}\Delta}{n - 1} \leq \delta \frac{V^*(\alpha)}{n - 1}, \quad (\text{B1})$$

for all $i \in N$, $p, p' \in \{R, S\} \times [0, 1] \times X$, and $\alpha \in A$.

To establish the result it suffices to show that every payoff vector in the Pareto frontier $W^* \equiv \{ (w_1, \dots, w_n) \in \mathbb{R}_+^n : \sum_{i=1}^n w_i = V^*(\alpha_0) \}$ can be supported in an equilibrium for all $\Delta < \bar{\Delta}$. Take an arbitrary $w^0 \in W^*$, and let $y^0 \in X$ be defined by $y_i^0 \equiv w_i^0 / V^*(\alpha_0)$ for all $i \in N$. Our objective is to construct a PBE σ , in which: (i) the committee implements the optimal stopping rule in each period; (ii) on the path, aggregate revenues are distributed according to y^0 in every period; and (iii) if committee member i deviates from σ in period t then, from $t + 1$ on, aggregate revenues are distributed according to $y^i \in X$, defined by

$$y_j^i \equiv \begin{cases} 0 & \text{if } j = i, \\ 1/(n - 1) & \text{if } j \neq i, \end{cases}$$

for all $j \in N$. As the optimal stopping rule is implemented in every continuation game both on and off the path, such an equilibrium must be renegotiation-proof.

More precisely, let

$$a^*(\alpha) \equiv \begin{cases} R & \text{if } \alpha > \alpha^* , \\ S & \text{otherwise.} \end{cases}$$

We define the strategy profile σ in terms of “phases,” formally represented by pairs in $\{1, \dots, n\} \times \{0, 1, \dots, n\}$. Every phase (ℓ, i) prescribes behavior in the ℓ th proposal stage of any given period and in the n voting stages that follow it: “ i ” indicates that σ prescribes policy $(a^*(\alpha), 1, y^i)$ to be implemented so that player i is punished. Specifically, in any period t in which the belief is $\alpha \in A$ and the order of proposers is $(\pi_1, \dots, \pi_n)$, if the game is in phase $(\ell, i) \neq (n, \pi_n)$, then σ prescribes the following behavior: (i) proposer π_ℓ offers policy $(a^*(\alpha), 1, y^i)$; (ii) if π_ℓ offered $(a^*(\alpha), 1, y^i)$, then every voter accepts it (irrespective of the previous voters’ behavior); and (iii) if π_ℓ offered any $p \neq (a^*(\alpha), 1, y^i)$, then every voter rejects it (irrespective of the previous voters’ behavior). If the game is in phase (n, π_n) , then σ prescribes the following behavior: (i) proposer π_n passes; and (ii) every voter rejects any proposal. Phases evolve according to the following recursive rules. In period 1, play begins in phase $(1, 0)$. Then in every period, at the end of any sequence of votes that began in any phase $(\ell, i) \neq (n, \pi_n)$:¹¹

- If policy $(a^*(\alpha), 1, y^i)$ was proposed and (if there was a vote) unanimously accepted, then the game transitions to phase $(1, i)$;
- if policy $(a^*(\alpha), 1, y^i)$ was accepted but not unanimously, then the game transitions to phase $(1, j)$, where j is the identity of the last voter who rejected $(a^*(\alpha), 1, y^i)$;
- if policy $(a^*(\alpha), 1, y^i)$ was proposed and rejected, then the game transitions to phase $(\ell + 1, j)$, where j is the identity of the last voter who rejected $(a^*(\alpha), 1, y^i)$;
- if the status quo differs from $(a^*(\alpha), 1, y^i)$ and proposer π_ℓ passes, then the game transitions to phase $(\ell + 1, \pi_\ell)$;
- if proposer π_ℓ offers to amend the status quo to a policy different from $(a^*(\alpha), 1, y^i)$ and her offer is unanimously rejected, then the game transitions to phase $(\ell + 1, \pi_\ell)$;

¹¹In what follows, we set $\ell + 1 = 1$ whenever $\ell = n$.

- if the proposer offers to amend the status quo to a policy different from $(a^*(\alpha), 1, y^i)$ and her offer is rejected but not unanimously, then the game transitions to phase $(\ell + 1, j)$, where j is the identity of the last voter who accepted the proposal.

- if the proposer offers to amend the status quo to a policy different from $(a^*(\alpha), 1, y^i)$ and her offer is accepted, then the game transitions to phase $(1, j)$, where j is the identity of the last voter who accepted the proposal.

At the end of any sequence of votes that began in phase (n, π_n) : if proposer π_n 's proposal is unanimously rejected, then the game transitions to phase $(1, \pi_n)$; otherwise, the game transitions to phase $(1, j)$, where j is the identity of the last voter who accepted the proposal.

We now verify that for $\Delta < \bar{\Delta}$, this strategy profile is a PBE. Observe that, by construction, each committee member j 's continuation value at the start of any phase $(\ell, i) \neq (n, \pi_n)$ is equal to

$$y_j^i V^*(\alpha) = \begin{cases} 0 & \text{if } j = i, \\ V^*(\alpha)/(n-1) & \text{if } j \neq i, \end{cases}$$

if that phase begins with belief α . As $\Delta < \bar{\Delta}$, it follows from (B1) that j 's objective is to avoid that the game transitions to phase j . More precisely, consider j 's voting behavior in phase $(\ell, i) \neq (n, \pi_n)$ when the proposer has offered policy $(a^*(\alpha), 1, y^i)$ different from the status quo. If all previous voters have accepted the proposal, or if j is the first voter, then $(a^*(\alpha), 1, y^i)$ will be implemented in the current period, irrespective of her decision. Her choice will only determine whether the game will transition to phase $(\ell + i, i)$ or to phase $(\ell + 1, j)$. As $y_j^i \geq y_j^j$ for all $i, j \in N$, she is better off accepting the proposal. Now suppose that at least one of the previous voters has rejected proposal $(a^*(\alpha), 1, y^i)$; let k be the identity of the last voter who rejected. In this case, committee member j 's vote will determine whether the game transitions to phase $(\ell + 1, k)$ (if she votes to accept) or to phase $(\ell + 1, j)$. It follows from (B1) and $\Delta < \bar{\Delta}$ that she is better off accepting, as prescribed by σ . The argument for the case where the proposer has offered to amend the status quo to a policy different from $(a^*(\alpha), 1, y^i)$ is analogous if either $j \neq \pi_n$ or $(\ell, i) \neq (n-1, \pi_n)$ (or both): voter j acts in accordance with σ to avoid a transition to

phase $(\ell + 1, j)$.

Consider voter π_n 's behavior in phase $(n - 1, \pi_n)$. Suppose first that proposer π_{n-1} has proposed policy $(a^*(\alpha), 1, y^{\pi_n})$. If all previous voters have accepted the proposal, or if π_n is the first voter, then $(a^*(\alpha), 1, y^{\pi_n})$ will be implemented in the current period and the game will transition to phase $(1, \pi_n)$, irrespective of π_n 's action. Hence, she cannot profitably deviate from σ in this case. If at least one of the previous voters has rejected proposal $(a^*(\alpha), 1, y^{\pi_n})$, then π_n 's payoff depends on her decision. Let α be the current belief; and let $p = (a, \tau, x)$ (resp. $p' = (a', \tau', x')$) be the policy that will be implemented in the current period if π_n votes to accept (resp. reject) the proposal — p and p' are determined by the current status quo policy and by the actions prescribed by σ to the remaining voters. By definition of σ , π_n is better off accepting if

$$(1 - \delta)w_{\pi_n}(p \mid \alpha) + \delta \frac{1}{n - 1} \mathbb{E}[V^*(\tilde{\alpha}) \mid \alpha, p] \geq (1 - \delta)w_{\pi_n}(p' \mid \alpha) ,$$

where $\mathbb{E}[V^*(\tilde{\alpha}) \mid \alpha, p]$ is the expected value of the belief at the start of the next period conditional on the current belief α and on p being implemented in the current period. It follows from (B1) and $\Delta < \bar{\Delta}$ that this inequality holds strictly and, therefore, that π_n cannot profitably deviate.

At the start of any phase (ℓ, i) with $\ell < n$ that begins with belief $\alpha \in A$, the proposer can either propose $(a^*(\alpha), 1, y^i)$ — in which case she receives payoff $y_{\pi_\ell}^i V^*(\alpha)$ — or propose any other policy — in which case she receives $y_{\pi_\ell}^{\pi_\ell} V^*(\alpha)$. As $y_{\pi_\ell}^i \geq y_{\pi_\ell}^{\pi_\ell}$, she is better off proposing $(a^*(\alpha), 1, y^i)$, as prescribed by σ . At the start of any phase (n, i) with $i \neq \pi_n$, proposer π_n can either propose $(a^*(\alpha), 1, y^i)$ — in which case she receives $y_{\pi_n}^i V^*(\alpha)$ — or propose any other policy — in which case she receives $(1 - \delta)w_{\pi_n}(p \mid \alpha) + \delta \mathbb{E}[V^*(\tilde{\alpha}) \mid \alpha, p]$, where p is the status quo. It follows from (B1) and $\Delta < \bar{\Delta}$ that she is strictly better off proposing $(a^*(\alpha), 1, y^i)$.

Consider now phase (n, π_n) , and let $\alpha \in A$ be the committee members' belief. The argument to show that no voter $j \neq \pi_n$ has a profitable deviation from σ in this phase is the same as above: as $\Delta < \bar{\Delta}$, j 's objective is to avoid a transition to phase j , irrespective of the policy implemented in the current period. As for voter π_n , her decision has no impact on her continuation payoff if all the previous voters rejected the proposal (as prescribed

by σ), and she is strictly better off inducing a transition to a phase $(1, k)$ with $k \neq \pi_n$ if some previous voter accepted the proposal. Finally, whether she passes (as prescribed by σ) or makes an unsuccessful proposal, proposer π_n 's payoff is the same, and therefore, any deviation is unprofitable.

C Proof of Proposition 3

Fix $\hat{\tau} > 0$. To prove the first part of Proposition 3, we will first define the threshold $\bar{\Delta}$ (Subsection C.1). Then, for every $\Delta < \bar{\Delta}$, we will construct a stationary Markov strategy profile σ^Δ that supports the optimal stopping rule (Subsection C.2). Finally, we will demonstrate that, for all $\Delta < \bar{\Delta}$, σ^Δ is a renegotiation-proof equilibrium (Subsection C.3).

C.1 Definition of $\bar{\Delta}$

To begin we must establish some notation. For each $i \in N$, let coalition C^i be defined by $C^i = N \setminus \{n\}$ if $i = 1$, and $C^i = N \setminus \{i - 1\}$ otherwise. Note that, as $\mathcal{D}$ is non-collegial, each coalition C^i is winning. Let $x^i \in X$ be defined by

$$x_j^i \equiv \begin{cases} 1/(n-1) & \text{if } j \in C^i, \\ 0 & \text{otherwise,} \end{cases}$$

and let

$$\bar{x}_i \equiv \frac{1}{n-1} \sum_{j: C^j \ni i} p_j,$$

where $p_j \in (0, 1)$ is the probability (induced by the protocol) that committee member j proposes first in any period. Next, let $k^* \in \mathbb{N}$ be implicitly defined by $\alpha_{k^*} \equiv \alpha^*$ and, for every $i, j \in N$, let the function $W_i^j: A \rightarrow \mathbb{R}_+$ be defined by

$$W_i^j(\alpha) \equiv \begin{cases} w_i(R, \hat{\tau}, x^j \mid 1) & \text{if } \alpha = 1, \\ w_i(S, \hat{\tau}, x^j \mid \alpha) & \text{if } \alpha = \alpha_k \text{ with } k \geq k^*, \\ [1 - \delta^{k^*-k}(1 - \gamma\Delta)^{k^*-k}]w_i(R, \hat{\tau}, x^j \mid \alpha) \\ + \delta^{k^*-k}[1 - \alpha_k + (1 - \gamma\Delta)^{k^*-k}\alpha_k]w_i(S, \hat{\tau}, x^j \mid \alpha) & \text{if } \alpha = \alpha_k \text{ with } k < k^*, \end{cases}$$

for all $\alpha \in A$. In words: for every α , $W_i^j(\alpha)$ is committee member i 's average discounted payoff is implemented along with (constant) redistributive policy $(\hat{\tau}, x^j)$ when the belief is α . Finally, let $W_i^0: A \rightarrow \mathbb{R}_+$ be defined by $W_i^0(\alpha) \equiv \sum_{j \in N} p_j W_i^j(\alpha)$, for all $\alpha \in A$. The interpretation of $W_i^0(\alpha)$ is analogous to $W_i^j(\alpha)$'s, but each redistributive policy $(\hat{\tau}, x^j)$ is implemented with probability p_j . Observe that, for all $\alpha \in A$ and $j \in N$, the payoff vectors $(W_i^j(\alpha))_{i \in N}$ and $(W_i^0(\alpha))_{i \in N}$ belong to the Pareto frontier. This observation will play an important role in the equilibrium construction below.

The definition of the threshold $\bar{\Delta}$ hinges on the following lemma, whose proof can be found in the supplementary appendix.

Lemma C1. *Suppose $\hat{\tau} > 0$ and, for all $i, j \in N$, let W_i^j and W_i^0 be defined as above. There exists $\bar{\Delta} > 0$ such that the following inequalities hold for all $\Delta < \bar{\Delta}$, all $i, j \in N$ with $i \in C^j$, and all $k \in \mathbb{N}$:*

$$\begin{aligned} W_i^j(1) &> (1 - \delta)\gamma\Delta\bar{r} + \delta W_i^0(1) , \\ W_i^j(\alpha_k) &> (1 - \delta)\bar{s}\Delta + \delta W_i^0(\alpha_k) , \text{ and} \\ W_i^j(\alpha_k) &> (1 - \delta)\alpha_k\gamma\Delta\bar{r} + \delta\alpha_k\gamma\Delta W_i^0(1) + \delta(1 - \alpha_k\gamma\Delta)W_i^0(\alpha_{k+1}) . \end{aligned}$$

Let $\bar{\Delta}$ be defined as in the above lemma. Henceforth, we assume that $\Delta < \bar{\Delta}$.

C.2 Definition of Stationary Markov Strategy Profile σ^Δ

This subsection describes the behavior prescribed by strategy profile σ^Δ to each committee member $i \in N$, for all $\Delta < \bar{\Delta}$. Observe that, in each proposal stage, i 's behavior only depends on the current status quo and belief and, in each voting stage, her behavior only depends on the current status quo, the belief, and the list of remaining proposers in the current period. Hence, σ^Δ is stationary Markov.

• **Proposal stages.** Consider first proposer i 's behavior in a period where the order of proposers is $\pi = (\pi_1, \dots, \pi_n)$ with $\pi_\ell = i$ for some ℓ ; and the first $\ell - 1$ proposers have failed to amend the status quo. There are three cases:

Case P1: The belief is α_k , where $k < k^*$.

Proposer i offers $(R, \hat{\tau}, x^i)$ (which, in cases where the status quo is $(R, \hat{\tau}, x^i)$ means that she passes).

Case P2: The belief is α_k , where $k > k^*$.

Proposer i offers $(S, \hat{\tau}, x^i)$ (which, in cases where the status quo is $(S, \hat{\tau}, x^i)$ means that she passes).

Case P3: The belief is α^* .

Case 3.1: If the status quo is a policy $(a, \tau, x) \neq (R, \hat{\tau}, x^j)$ for all $j \in N$, then proposer i offers $(S, \hat{\tau}, x^i)$.

Case 3.2: If the status quo is $(R, \hat{\tau}, x^j)$ for some $j \in N$, then proposer i offers $(S, \hat{\tau}, x^j)$ if $i \in C^j$, and $(S, \hat{\tau}, x^i)$ otherwise.

• **Voting stages.** Consider now voter i 's behavior in a period where the order of proposers is $\pi = (\pi_1, \dots, \pi_n)$. There are several cases:

Case V1: The status quo is $(R, \hat{\tau}, x^j)$ for some $j \in N$; the belief is α_k , where $k < k^*$; and a proposer π_ℓ has just proposed policy $(a, \tau, y) \neq (R, \hat{\tau}, x^j)$.

If voter i is a member of C^j , then she rejects the proposal; otherwise, she accepts the proposal if and only if

$$W_i^j(\alpha_k) > (1-\delta)w_i(a, \tau, y \mid \alpha_k) + \delta \begin{cases} \alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k \gamma \Delta) W_i^0(\alpha_{k+1}) & \text{if } a = R, \\ W_i^0(\alpha_k) & \text{if } a = S. \end{cases}$$

Case V2: The status quo is (a, τ, x) , where $(a, \tau, x) \neq (R, \hat{\tau}, x^j)$ for all $j \in N$; the belief is α_k , where $k < k^*$.

Case V2.1: Proposer π_n has just proposed policy $(R, \hat{\tau}, x^j)$ for some $j \in N$.

If voter i is a member of C^j , then she accepts the proposal; otherwise, she accepts the proposal if and only if

$$W_i^j(\alpha_k) > (1-\delta)w_i(a, \tau, x \mid \alpha_k) + \delta \begin{cases} \alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k \gamma \Delta) W_i^0(\alpha_{k+1}) & \text{if } a = R, \\ W_i^0(\alpha_k) & \text{if } a = S. \end{cases}$$

Case V2.2: Proposer π_n has just proposed policy (b, τ', y) , where $(b, \tau', y) \neq (R, \hat{\tau}, x^j)$ for all $j \in N$.

Voter i accepts the proposal if and only if

$$(1-\delta)w_i(a, \tau, x \mid \alpha_k) + \delta \begin{cases} \alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k \gamma \Delta) W_i^0(\alpha_{k+1}) & \text{if } a = R, \\ W_i^0(\alpha_k) & \text{if } a = S. \end{cases}$$

$$< (1-\delta)w_i(b, \tau', y \mid \alpha_k) + \delta \begin{cases} \alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k \gamma \Delta) W_i^0(\alpha_{k+1}) & \text{if } b = R, \\ W_i^0(\alpha_k) & \text{if } b = S. \end{cases}$$

Case V2.3: Proposer π_ℓ , $\ell < m$, has just proposed policy $(R, \hat{\tau}, x^j)$ for some $j \in N$.

If voter i is a member of C^j , then she accepts the proposal; otherwise, she accepts the proposal if and only if $W_i^j(\alpha_k) \geq W_i^{\pi_{\ell+1}}(\alpha_k)$.

Case V2.4: Proposer π_ℓ , $\ell < m$, has just proposed policy (b, τ', y) , where $(b, \tau', y) \neq (R, \hat{\tau}, x^j)$ for all $j \in N$.

Voter i accepts the proposal if and only if

$$W_i^{\pi_{\ell+1}}(\alpha_k) < (1-\delta)w_i(b, \tau', y \mid \alpha_k) + \delta \begin{cases} \alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k \gamma \Delta) W_i^0(\alpha_{k+1}) & \text{if } b = R, \\ W_i^0(\alpha_k) & \text{if } b = S. \end{cases}$$

Case V3: The status quo is $(S, \hat{\tau}, x^j)$ for some $j \in N$; the belief is α_k , where $k > k^*$; and a proposer π_ℓ has just proposed policy $(a, \tau, y) \neq (S, \hat{\tau}, x^j)$.

If voter i is a member of C^j , then she rejects the proposal; otherwise, she accepts the proposal if and only if

$$W_i^j(\alpha_k) > (1-\delta)w_i(a, \tau, y \mid \alpha_k) + \delta \begin{cases} \alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k \gamma \Delta) W_i^0(\alpha_{k+1}) & \text{if } a = R, \\ W_i^0(\alpha_k) & \text{if } a = S. \end{cases}$$

Case V4: The status quo is (a, τ, x) , where $(a, \tau, x) \neq (S, \hat{\tau}, x^j)$ for all $j \in N$; and the belief is α_k , where $k > k^*$.

Case V4.1: Proposer π_n has just proposed policy $(S, \hat{\tau}, x^j)$ for some $j \in N$.

If voter i is a member of C^j , then she accepts the proposal; otherwise, she accepts the proposal if and only if

$$W_i^j(\alpha_k) \geq (1-\delta)w_i(a, \tau, x \mid \alpha_k) + \delta \begin{cases} \alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k \gamma \Delta) W_i^0(\alpha_{k+1}) & \text{if } a = R, \\ W_i^0(\alpha_k) & \text{if } a = S. \end{cases}$$

Case V4.2: Proposer π_n has just proposed policy (b, τ', y) , where $(b, \tau', y) \neq (S, \hat{\tau}, x^j)$ for all $j \in N$.

Voter i accepts the proposal if and only if

$$\begin{aligned} & (1-\delta)w_i(a, \tau, x \mid \alpha_k) + \delta \begin{cases} \alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k \gamma \Delta) W_i^0(\alpha_{k+1}) & \text{if } a = R, \\ W_i^0(\alpha_k) & \text{if } a = S. \end{cases} \\ & \leq (1-\delta)w_i(b, \tau', y \mid \alpha_k) + \delta \begin{cases} \alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k \gamma \Delta) W_i^0(\alpha_{k+1}) & \text{if } b = R, \\ W_i^0(\alpha_k) & \text{if } b = S. \end{cases} \end{aligned}$$

Case V4.3: Proposer π_ℓ , $\ell < m$, has just proposed policy $(S, \hat{\tau}, x^j)$ for some $j \in N$.

If voter i is a member of C^j , then she accepts the proposal; otherwise, she accepts the proposal if and only if $W_i^j(\alpha_k) \geq W_i^{\pi_{\ell+1}}(\alpha_k)$.

Case V4.4: Proposer π_ℓ , $\ell < m$, has just proposed policy (b, τ', y) , where $(b, \tau', y) \neq (S, \hat{\tau}, x^j)$ for all $j \in N$.

Voter i accepts the proposal if and only if

$$W_i^{\pi_{\ell+1}}(\alpha_k) < (1-\delta)w_i(b, \tau', y \mid \alpha_k) + \delta \begin{cases} \alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k \gamma \Delta) W_i^0(\alpha_{k+1}) & \text{if } b = R, \\ W_i^0(\alpha_k) & \text{if } b = S. \end{cases}$$

Case V5: The status quo is $(R, \hat{\tau}, x^j)$ for some $j \in N$; and the belief is α^* .

Case V5.1: Proposer π_n has just proposed policy $(S, \hat{\tau}, x^{j'})$ for some $j' \in N$.

If voter i is a member of $C^{j'}$, then she accepts the proposal; otherwise, she accepts the proposal if and only if

$$W_i^{j'}(\alpha^*) \geq (1-\delta)w_i(R, \hat{\tau}, x^j \mid \alpha^*) + \delta \alpha^* \gamma \Delta W_i^0(1) + \delta (1 - \alpha^* \gamma \Delta) W_i^0(\alpha_{k^*+1}) .$$

Case V5.2: Proposer π_n has just proposed policy (b, τ', y) , where $(b, \tau', y) \neq (S, \hat{\tau}, x^j)$ for all $j \in N$.

Voter i accepts the proposal if and only if

$$(1 - \delta)w_i(R, \hat{\tau}, x^j \mid \alpha^*) + \delta\alpha^*\gamma\Delta W_i^0(1) + \delta(1 - \alpha^*\gamma\Delta)W_i^0(\alpha_{k^*+1}) \\ \leq (1 - \delta)w_i(b, \tau', y \mid \alpha^*) + \delta \begin{cases} \alpha^*\gamma\Delta W_i^0(1) + (1 - \alpha^*\gamma\Delta)W_i^0(\alpha_{k^*+1}) & \text{if } b = R, \\ W_i^0(\alpha^*) & \text{if } b = S. \end{cases}$$

Case V5.3: Proposer π_ℓ , $\ell < m$, has just proposed policy $(S, \hat{\tau}, x^{j'})$ for some $j' \in N$.

If voter i is a member of C^j and $j' = j$, then she accepts the proposal; if she is a member of C^j , $j' \neq j$ and there is $\ell' > \ell$ such that $\pi_{\ell'} \in C^j$, then she rejects the proposal; otherwise, she accepts the proposal if and only if $W_i^{j'}(\alpha^*) > W_i^{\pi_{\ell+1}}(\alpha^*)$.

Case V5.4: Proposer π_ℓ , $\ell < m$, has just proposed policy (b, τ', y) , where $(b, \tau', y) \neq (R, \hat{\tau}, x^j)$ for all $j \in N$.

Voter i accepts the proposal if and only if

$$W_i^{\pi_{\ell+1}}(\alpha^*) < (1 - \delta)w_i(b, \tau', y \mid \alpha^*) + \delta \begin{cases} \alpha^*\gamma\Delta W_i^0(1) + (1 - \alpha^*\gamma\Delta)W_i^0(\alpha_{k^*+1}) & \text{if } b = R, \\ W_i^0(\alpha^*) & \text{if } b = S. \end{cases}$$

Case V6: The status quo is (a, τ, x) , with $(a, \tau, x) \neq (R, \hat{\tau}, x^j)$ for all $j \in N$; and the belief is α^* .

Voter i behaves as in Cases 3 and 4 (with $k = k^*$).

Case V7: The belief is equal to one. In this case, apply Cases V1 and V2 replacing α_k and α_{k+1} by 1.

C.3 Verification that σ^Δ Is a Renegotiation-proof equilibrium

Optimal stopping rule. Before proceeding to the verification that σ^Δ is a renegotiation-proof equilibrium, it is worth noting that it sustains the optimal stopping rule. The game starts with status quo $(S, 0, x^0)$ and belief α_0 . The first proposer, say j , is prescribed

to propose policy $(R, \hat{\tau}, x^j)$ (see Case P1 above), which is accepted by all the members of decisive coalition C^j (see Case V2.3). This policy is then implemented again in every future period that begins with a belief greater than α^* , as any proposal to amend it is voted down by the members of C^j (see Case V1). If the belief becomes α^* , then all members of C^j reject any proposal until one of them proposes policy $(S, \hat{\tau}, x^j)$, which they all accept (see Cases P3.2 and V5.3). (As C^j is a decisive coalition, $C^j \cap \{\pi_1, \dots, \pi_m\} \neq \emptyset$ and, therefore, at least one of its members is a proposer.) Policy $(S, \hat{\tau}, x^j)$ is then never amended, as any proposal to change it is voted down by the members of C^j (see Case V3). Hence, the optimal stopping rule is implemented on the path. Any deviation from this path leads the next proposer j' to successfully propose policy $(R, \hat{\tau}, x^{j'})$ if the belief is greater than α^* , or policy $(S, \hat{\tau}, x^{j'})$ if the belief smaller than or equal to α^* . The induced path again supports the optimal stopping rule.

Continuation values and renegotiation-proofness. Let $V_i(b, \tau, y|\alpha)$ be player i 's average discounted value (induced by σ^Δ) from implementing policy (b, τ, y) when the belief is equal to α . Observe first that if a policy $(R, \hat{\tau}, x^j)$, with $j \in N$, is implemented in a period that begins with belief α_k , $k < k^*$, then it is also implemented in any future period beginning with a belief greater than α^* (Cases V1 and V7 above). If the belief becomes equal to α^* , then $(R, \hat{\tau}, x^j)$ is amended to policy $(S, \hat{\tau}, x^j)$ (see Cases P3.2, V5.1 and V5.3), which is then implemented in every future period (see Cases V6 and V3). This implies that $V_i(R, \hat{\tau}, x^j|\alpha_k) = W_i^j(\alpha_k)$ for all $i, j \in N$ and $k < k^*$. Similar arguments establish that $V_i(R, \hat{\tau}, x^j|1) = W_i^j(1)$ and $V_i(S, \hat{\tau}, x^j|\alpha_k) = W_i^j(\alpha_k)$, for all $i, j \in N$ and $k \geq k^*$. By construction of σ^Δ , these are all the possible continuation values induced by σ^Δ at the start of any continuation game. As they all belong to the Pareto frontier (Subsection C.1), this implies that if σ^Δ is an equilibrium, then it must be renegotiation proof.

Now suppose that a policy (b, τ, y) , where $(b, \tau, y) \neq (R, \hat{\tau}, x^j)$ for all $j \in N$, is implemented in a period that starts with a belief α_k , $k < k^*$. Player i receives $(1 - \delta)w_i(R, \tau, y | \alpha_k)$ in that period. If b is alternative R , then there are two possible cases: it succeeds with probability $\alpha_k \gamma \Delta$, in which case the next period's first proposer j successfully offers policy

$(R, \hat{\tau}, x^j)$ (see Cases P1, V7 and V1); and it fails with probability $1 - \alpha_k \gamma \Delta$ in which case, the next period's first proposer j successfully offers policy $(R, \hat{\tau}, x^j)$ if $k < k^* - 1$ (see Case V2.3), or $(S, \hat{\tau}, x^j)$ if $k = k^* - 1$ (see Cases V6 and V4.3). Therefore,

$$\begin{aligned}
V_i(R, \tau, y | \alpha_k) &= (1 - \delta)w_i(R, \tau, y | \alpha_k) + \delta \alpha_k \gamma \Delta \sum_{j \in N} p_j V_i(R, \hat{\tau}, x^j | 1) \\
&\quad + \delta(1 - \alpha_k \gamma \Delta) \sum_{j \in N} p_j \begin{cases} V_i(R, \hat{\tau}, x^j | \alpha_{k+1}) & \text{if } k < k^* - 1 \\ V_i(S, \hat{\tau}, x^j | \alpha^*) & \text{if } k = k^* - 1 \end{cases} \\
&= (1 - \delta)w_i(b, \tau, y | \alpha_k) + \delta \alpha_k \gamma \Delta \sum_{j \in N} p_j W_i^j(1) \\
&\quad + \delta(1 - \alpha_k \gamma \Delta) \sum_{j \in N} p_j \begin{cases} W_i^j(\alpha_{k+1}) & \text{if } k < k^* - 1 \\ W_i^j(\alpha^*) & \text{if } k = k^* - 1 \end{cases} \\
&= (1 - \delta)w_i(R, \tau, y | \alpha_k) + \delta [\alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k) \gamma \Delta W_i^0(\alpha_{k+1})] .
\end{aligned}$$

If b is alternative S then, in the next period, the first proposer j successfully offers policy $(R, \hat{\tau}, x^j)$ (see Cases P1 and V2.3); so that

$$\begin{aligned}
V_i(S, y | \alpha_k) &= (1 - \delta)w_i(S, \tau, y | \alpha_k) + \delta \sum_{j \in M} p_j V_i(s, \hat{\tau}, x^j | \alpha_k) \\
&= (1 - \delta)w_i(S, \tau, y | \alpha_k) + \delta \sum_{j \in N} p_j W_i^j(1) \\
&= (1 - \delta)w_i(S, \tau, y | \alpha_k) + \delta W_i^0(1) .
\end{aligned}$$

Using parallel arguments, one can show that player i 's continuation value from implementing a policy (b, τ, y) , where $(b, \tau, y) \neq (S, \hat{\tau}, x^j)$ for all $j \in N$, in a period with belief α_k , $k \geq k^*$, is given by

$$V_i(b, \tau, y | \alpha_k) = (1 - \delta)w_i(b, \tau, y | \alpha_k) + \delta \begin{cases} W_i^0(\alpha_{k+1}) & \text{if } b = r , \\ W_i^0(\alpha_k) & \text{if } b = s , \end{cases}$$

and that her continuation value from implementing a policy (b, τ, y) , where $(b, \tau, y) \neq (R, \hat{\tau}, x^j)$ for all $j \in N$, in a period where the belief is equal to one is given by

$$V_i(b, \tau, y | 1) = (1 - \delta)w_i(b, \tau, y | 1) + \delta W_i^0(1) .$$

It follows directly from this characterization of continuation values and from Lemma C1 that $(R, \hat{\tau}, x^j)$ with $i \in C^j$ is player i 's ideal policy when the belief is greater than α^* — i.e., $V_i(R, \hat{\tau}, x^j | \alpha) \geq V_i(b, \tau, y | \alpha)$ for all $i, j \in N$ such that $i \in C^j$, $(b, \tau, y) \in \{R, S\} \times [0, \hat{\tau}] \times X$, and $\alpha > \alpha^*$ — and that $(S, \hat{\tau}, x^j)$ with $i \in C^j$ is her ideal policy when the belief is smaller than or equal to α^* — i.e., $V_i(S, \hat{\tau}, x^j | \alpha) \geq V_i(b, \tau, y | \alpha)$ for all $i, j \in N$ such that $i \in C^j$, $(b, \tau, y) \in \{R, S\} \times [0, \hat{\tau}] \times X$, and $\alpha \leq \alpha^*$.

Voting stages. To verify that σ^Δ is an equilibrium, we will first check that all possible deviations in voting stages are unprofitable. To do so, we will consider in turn the various cases in the definition of voting strategies.

In Case V1, (decisive) voter i receives a payoff of $V_i(R, \hat{\tau}, x^j | \alpha_k) = W_i^j(\alpha_k)$ if she rejects the proposal (a, τ, y) to amend status quo $(R, \hat{\tau}, x^j)$ (as any future attempt to amend it in this period will be unsuccessful), and a payoff of

$$V_i(a, \tau, y | \alpha_k) = (1-\delta)w_i(a, \tau, y | \alpha_k) + \delta \begin{cases} \alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k \gamma \Delta) W_i^0(\alpha_{k+1}) & \text{if } a = R, \\ W_i^0(\alpha_k) & \text{if } a = S, \end{cases}$$

if she accepts it. Hence, she cannot profitably deviate from σ^Δ if she is not a member of C^j . Moreover, it follows from Lemma C1 and the above equality that $V_i(a, \tau, y | \alpha_k) < W_i^j(\alpha_k) = V_i(R, \hat{\tau}, x^j | \alpha_k)$ for all $i \in C^j$, so that voter i cannot profitably deviate from rejecting (a, τ, y) if she is a member of C^j .

It follows immediately from our characterization of continuation values above that, in Cases V2.1 and V2.2, σ^Δ prescribes voter i to accept the last proposer's offer if and only if her continuation value from implementing this offer exceeds her continuation value from implementing the status quo. Therefore, deviations are also unprofitable in these cases. In case V2.3, (decisive) player i anticipates that if she rejects the ℓ th proposer's offer, $(R, \hat{\tau}, x^j)$, then the next proposer will successfully propose $(R, \hat{\tau}, x^{\pi_{\ell+1}})$. It is therefore optimal for her to accept $(R, \hat{\tau}, x^j)$ if and only if $V_i(R, \hat{\tau}, x^j | \alpha_k) = W_i^j(\alpha_k) \geq W_i^{\pi_{\ell+1}}(\alpha_k) = V_i(R, \hat{\tau}, x^{\pi_{\ell+1}} | \alpha_k)$, as prescribed by σ^Δ to every $i \notin C^j$. Moreover, by definition of W_i^j , $W_i^j(\alpha_k) \geq W_i^{j'}(\alpha_k)$ for all $i, j, j' \in N$ such that $i \in C^j$. Therefore, it is always optimal for voter i to accept $(R, \hat{\tau}, x^j)$ if $i \in C^j$. The same argument applies to Case 2.4 except that

in this case, the ℓ th proposal offer is some $(b, \tau, y) \neq (R, \hat{\tau}, x^j)$ for all $j \in N$, so that the value from accepting it is equal to

$$V_i(b, \tau, y) = (1 - \delta)w_i(b, \tau, y \mid \alpha_k) + \delta \begin{cases} \alpha_k \gamma \Delta W_i^0(1) + (1 - \alpha_k \gamma \Delta) W_i^0(\alpha_{k+1}) & \text{if } b = R, \\ W_i^0(\alpha_k) & \text{if } b = S. \end{cases}$$

The arguments to show that there is no profitable deviation from σ^Δ in all the other cases, but Case V5.3, are analogous: in each of these cases, σ^Δ prescribes decisive voter i the action that maximizes her continuation value. In Case V5.3, (decisive) voter i anticipates that if she rejects the ℓ th proposer's offer, $(R, \hat{\tau}, x^{j'})$, then the next proposer will successfully propose $(R, \hat{\tau}, x^{\pi_{\ell+1}})$, and she will consequently receive $V_i(R, \hat{\tau}, x^{\pi_{\ell+1}} \mid \alpha^*) = W_i^{\pi_{\ell+1}}(\alpha^*)$. If she is a member of coalition C^j and $j' = j$, then it is optimal for her to accept the offer (as prescribed by σ^Δ): by definition of W_i , $V_i(R, \hat{\tau}, x^{j'} \mid \alpha^*) = V_i(R, \hat{\tau}, x^j \mid \alpha^*) = W_i^j(\alpha^*) \geq W_i^{\pi_{\ell+1}}(\alpha^*)$. This implies that every member of coalition C^j knows that if the status quo is $(R, \hat{\tau}, x^j)$ in a period where the belief is α^* , then the first proposer in C^j will successfully offer policy $(S, \hat{\tau}, x^j)$. It therefore follows from Lemma C1 and our the characterization of continuation values above that every member i of C^j obtains her highest possible continuation value, $V_i(S, \hat{\tau}, x^j \mid \alpha^*) = W_i^j(\alpha^*)$, by rejecting any offer until a proposer in C^j successfully offers $(S, \hat{\tau}, x^j)$ (as prescribed by σ^Δ). Finally, if i is not a member of C^j , or if (off the path) all the proposers in C^j have failed to amend the status quo, then σ^Δ optimally prescribes her, as in the previous cases, to choose the action that maximizes her continuation value.

Proposal stages. Consider now player i 's behavior in a period where she is the ℓ th proposer and the first $\ell - 1$ proposers have failed to amend the status quo. We begin with cases where the belief α is greater than α^* . If the status quo is $(R, \hat{\tau}, x^j)$ for some $j \in N$, then any proposal to amend it is voted down by decisive coalition C^j (see Cases V1 and V7). Therefore, proposer i 's payoff will be $V_i(R, \hat{\tau}, x^j \mid \alpha)$, irrespective of the action she takes. If the status quo is a policy $(a, \tau, x) \neq (R, \hat{\tau}, x^j)$ for all $j \in N$ and she proposes $(R, \hat{\tau}, x^i)$ (as prescribed by σ^Δ), then her proposal is accepted by all the members of C^i (see Cases V2.1 and V2.3) and she obtains a payoff of $V_i(R, \hat{\tau}, x^i \mid \alpha) = W_i^i(\alpha)$. As we established

above, this is the highest payoff she can get if $\alpha > \alpha^*$. Therefore, any deviation from σ^Δ must be unprofitable.

Suppose now that the belief α is smaller than α^* . If the status quo is $(S, \hat{\tau}, x^j)$ for some $j \in N$, then any proposal to amend it is voted down by decisive coalition C^j (see Case V3). Therefore, proposer i 's payoff will be $V_i(S, \hat{\tau}, x^j | \alpha)$, irrespective of the action she takes. If the status quo is a policy $(a, \tau, x) \neq (S, \hat{\tau}, x^j)$ for all $j \in N$ and she proposes $(S, \hat{\tau}, x^i)$ (as prescribed by σ^Δ), then her proposal is accepted by all the members of C^i (see Cases V4.1 and V4.3) and she obtains a payoff of $V_i(S, \hat{\tau}, x^i | \alpha) = W_i^i(\alpha)$. As we established above, this is the highest payoff she can get if $\alpha \leq \alpha^*$. Therefore, any deviation from σ^Δ must again be unprofitable.

Finally, suppose that the belief is equal to α^* . If the status quo is a policy $(a, \tau, x) \neq (R, \hat{\tau}, x^j)$ for all $j \in N$, then the proof that proposer i cannot deviate from proposing $(S, \hat{\tau}, x^i)$ is the same as in the previous paragraph. If the status quo is a policy $(R, \hat{\tau}, x^j)$ for some $j \in N$, then there are several possible cases:

(i) If i is a member of C^j and she proposes $(R, \hat{\tau}, x^j)$, then her proposal is accepted by all the the members of decisive coalition C^j . As this policy is one of her ideal policies when the belief is α^* , she cannot profitably deviate from the behavior prescribed in Case P3.2.

(ii) If i is not a member of C^j and all proposers in C^j have already (unsuccessfully) proposed, then all the members of decisive coalition C^i accept proposal $(S, \hat{\tau}, x^i)$ (see Case V5.1 and the last sub-case in Case V5.3). As we established above, this is the policy that maximizes her continuation value when the belief is α^* . It is therefore impossible for her to profitably deviate from proposing it, as prescribed in Case P3.1.

(iii) If i is not a member of C^j and some proposers in C^j have not yet proposed, then any proposal that differs from $(S, \hat{\tau}, x^j)$ is rejected by the members of C^j (second sub-case in Case V5.3) and, by the end of the period, some proposer in C^j will successfully propose $(S, \hat{\tau}, x^j)$. This implies that, irrespective of the proposal she makes, proposer i 's payoff will be $V_i(S, \hat{\tau}, x^j | \alpha^*)$. Therefore, any deviation is unprofitable.

This proves that σ^Δ is a renegotiation-proof stationary equilibrium, thus completing

the proof of Proposition 3.

References

- Acemoglu, D., Golosov, M. and Tsyvinski, A. (2008), ‘Political Economy of Mechanisms’, *Econometrica* **76**(3), 619–641.
- Acemoglu, D., Naidu, S., Restrepo, P. and Robinson, J. A. (2015), *Democracy, Redistribution, and Inequality*, Vol. 2B, North Holland, chapter 21.
- Alesina, A. and Drazen, A. (1991), ‘Why Are Stabilizations Delayed?’, *American Economic Review* **81**(5), 1170–1188.
- Anesi, V. and Duggan, J. (2018), ‘Existence and Indeterminacy of Markovian Equilibria in Dynamic Bargaining Games’, *Theoretical Economics* **13**(2), 505–525.
- Anesi, V. and Seidmann, D. J. (2015), ‘Bargaining in Standing Committees with an Endogenous Default’, *The Review of Economic Studies* **82**(3), 825–867.
- Austen-Smith, D. and Banks, J. S. (1999), *Positive Political Theory I: Collective Preference*, University of Michigan Press.
- Baron, D. P. (1996), ‘A Dynamic Theory of Collective Goods Programs’, *The American Political Science Review* **90**(2), 316–330.
- Baron, D. P. and Bowen, T. R. (2015), ‘Dynamic Coalitions’, *Working Paper*.
- Battaglini, M. and Palfrey, T. R. (2012), ‘The Dynamics of Distributive Politics’, *Economic Theory* **49**(3), 739–777.
- Bergin, J. and MacLeod, W. B. (1993), ‘Continuous Time Repeated Games’, *International Economic Review* pp. 21–37.
- Besley, T. and Coate, S. (1998), ‘Sources of Inefficiency in a Representative Democracy: A Dynamic Analysis’, *American Economic Review* **88**(1), 139–156.

- Bouton, L., Llorente-Saguer, A. and Malherbe, F. (2018), ‘Get Rid of Unanimity Rule: The Superiority of Majority Rules with Veto Power’, *Journal of Political Economy* **126**(1), 107–149.
- Bowen, T. R., Chen, Y. and Eraslan, H. (2014), ‘Mandatory Versus Discretionary Spending: The Status Quo Effect’, *American Economic Review* **104**(10), 2941–2974.
- Bowen, T. R. and Zahran, Z. (2012), ‘On Dynamic Compromise’, *Games and Economic Behavior* **76**, 391–419.
- Cai, H. and Treisman, D. (2009), ‘Political Decentralization and Policy Experimentation’, *Quarterly Journal of Political Science* **4**(1), 35–58.
- Callander, S. (2011), ‘Searching for Good Policies’, *American Political Science Review* **105**, 643–662.
- Callander, S. and Hummel, P. (2014), ‘Preemptive Policy Experimentation’, *Econometrica* **82**(4), 1509–1528.
- Cho, S.-j., Duggan, J. et al. (2009), ‘Bargaining Foundations of the Median Voter Theorem’, *Journal of Economic Theory* **144**(2), 851–868.
- Dewatripont, M. and Roland, G. R. (1992), ‘Economic Reform and Dynamic Political Constraints’, *Review of Economic Studies* **59**, 703–730.
- Diermeier, D., Egorov, G. and Sonin, K. (2017), ‘Political Economy of Redistribution’, *Econometrica* .
- Dixit, A. and Olson, M. (2000), ‘Does Voluntary Participation Undermine the Coase Theorem?’, *Journal of Public Economics* **76**(3), 309–335.
- Dziuda, W. and Loeper, A. (2016), ‘Dynamic Collective Choice with Endogenous Status Quo’, *Journal of Political Economy* **124**(4).
- Dziuda, W. and Loeper, A. (2018), ‘Dynamic Pivotal Politics’, *American Political Science Review* pp. 1–22.

- Eraslan, H. and Merlo, A. (2002), ‘Majority Rule in a Stochastic Model of Bargaining’, *Journal of Economic Theory* **103**, 103.
- Fernandez, R. and Rodrik, D. (1991), ‘Resistance to Reform: Status Quo Bias in the Presence of Individual-Specific Uncertainty’, *American Economic Review* **81**(5), 1146–1155.
- Freer, M., Martinelli, C. and Wang, S. (2018), ‘Collective Experimentation: A Laboratory Study’, *Working Paper* .
- Hirsch, A. V. (2016), ‘Experimentation and Persuasion in Political Organizations’, *American Political Science Review* **110**(1), 68–84.
- Hirsch, A. V., Shotts, K. W. et al. (2015), ‘Veto Players and Policy Entrepreneurship’, *Working Paper* .
- Kalandrakis, T. (2004), ‘A Three-Player Dynamic Majoritarian Bargaining Game’, *Journal of Economic Theory* **116**(2), 294–322.
- Kalandrakis, T. (2010), ‘Minimum winning coalitions and endogenous status quo’, *International Journal of Game Theory* **39**, 617–643.
- Keller, G., Rady, S. and Cripps, M. (2005), ‘Strategic Experimentation with Exponential Bandits’, *Econometrica* **73**(1), 39–68.
- Krehbiel, K. (1998), *Pivotal Politics: A Theory of US Lawmaking*, University of Chicago Press.
- Maggi, G. and Morelli, M. (2006), ‘Self-Enforcing Voting in International Organizations’, *American Economic Review* **96**(4), 1137–1158.
- Majumdar, S. and Mukand, S. W. (2004), ‘Policy Gambles’, *The American Economic Review* **94**(4), 1207–1222.
- McCarty, N. (2000), ‘Proposal Rights, Veto Rights, and Political Bargaining’, *American Journal of Political Science* pp. 506–522.

- Merlo, A. and Wilson, C. (1995), ‘A Stochastic Model of Sequential Bargaining with Complete Information’, *Econometrica* pp. 371–399.
- Messner, M. and Polborn, M. K. (2012), ‘The Option to Wait in Collective Decisions and Optimal Majority Rules’, *Journal of Public Economics* **96**(5-6), 524–540.
- Murto, P. and Välimäki, J. (2011), ‘Learning and Information Aggregation in an Exit Game’, *The Review of Economic Studies* **78**(4), 1426–1461.
- Nunnari, S. (2014), ‘Dynamic Legislative Bargaining with Veto Power: Theory and Experiments’, *Working Paper* .
- Peritz, L. (2017), ‘When are International Institutions Effective? The Impact of Domestic Veto Players on Compliance with WTO Rulings’, *Working Paper* .
- Piguillem, F. and Riboni, A. (2015), ‘Spending-Biased Legislators: Discipline through Disagreement’, *The Quarterly Journal of Economics* **130**(2), 901–949.
- Richter, M. (2014), ‘Fully Absorbing Dynamic Compromise’, *Journal of Economic Theory* **152**, 92–104.
- Rubinstein, A. (1982), ‘Perfect Equilibrium in a Bargaining Model’, *Econometrica* pp. 97–109.
- Schelling, T. C. (1960), ‘The Strategy of Conflict’, *Cambridge, Mass* .
- Shaked, A. and Sutton, J. (1984), ‘Involuntary Unemployment as a Perfect Equilibrium in a Bargaining Model’, *Econometrica* pp. 1351–1364.
- Strulovici, B. (2010), ‘Learning while Voting: Determinants of Collective Experimentation’, *Econometrica* **78**(3), 933–971.
- Tornell, A. (1998), ‘Reform from Within’, *NBER Working Paper No. 6497* .
- Tsebelis, G. (2002), *Veto Players: How Political Institutions Work*, Princeton University Press.

Volden, C., Ting, M. M. and Carpenter, D. P. (2008), 'A Formal Model of Learning and Policy Diffusion', *The American Political Science Review* **102**(3), 319–332.

SUPPLEMENTARY APPENDIX: PROOFS OMITTED FROM TEXT

A Proof of Lemma A1

We begin with part (i). Suppose first that $L \notin \mathcal{D}$. To see that there exists an equilibrium with the proposed outcome, consider the following (stationary Markov) strategy profile:

- Whenever the status quo is R , all proposers pass (i.e., propose R), and each voter i accepts proposal S if and only if $i \in L$;
- whenever the status quo is S , each proposer i proposes R if $i \notin L$ and passes otherwise, and each voter i accepts proposal R if and only if $i \notin L$.

It is easy to check that this strategy profile constitutes an equilibrium. (In particular, proposers who prefer S to R do not deviate and propose to amend status quo R because they anticipate that such a proposal would be rejected.)

Next we show that this is the unique equilibrium outcome. Our proof shares some of the intuitions of the Shaked and Sutton (1984) proof of equilibrium uniqueness for the Rubinstein (1982) model. Let the set of PBEs of $\Gamma(R, 1)$ be denoted by $\mathcal{E}(R, 1)$. In $\Gamma(R, 1)$, committee member i 's expected payoff in every period t is a convex combination of $\gamma\Delta r_i$ and Δs_i . Therefore, for every strategy profile σ her average discounted payoff is of the form $V_i(\sigma) \equiv \beta(\sigma)\gamma\Delta r_i + [1 - \beta(\sigma)]\Delta s_i$, with $\beta(\sigma) \in [0, 1]$. This implies that, for any two strategy profiles σ and σ' , and any committee member $i \in \{i \in N : \gamma r_i > s_i\}$, we have $V_i(\sigma) \geq V_i(\sigma')$ if and only if $\beta(\sigma) \geq \beta(\sigma')$.

Let $\{\sigma^m\}$ be a sequence in $\mathcal{E}(R, 1)$ that satisfies $\lim_{m \rightarrow \infty} \beta(\sigma^m) = \inf_{\sigma \in \mathcal{E}(R, 1)} \beta(\sigma)$, so that $\lim_{m \rightarrow \infty} V_i(\sigma^m) = \inf_{\sigma \in \mathcal{E}(R, 1)} V_i(\sigma)$ for all $i \in W \equiv \{i \in N : \gamma r_i \geq s_i\}$. Fix $m \in \mathbb{N}$. Every proposal that may successfully be made by the last proposer in the first period under σ^m (both on and off the path) must be accepted by some decisive player i in W . That is, i 's continuation payoff from accepting the proposal, say U_i^a , must be

at least as large as her payoff from rejecting it; i.e., $U_i^a \geq (1 - \delta)\gamma\Delta r_i + \delta V_i(\sigma^r)$, where $\sigma^r \in \mathcal{E}(R, 1)$ is the equilibrium of $\Gamma(R, 1)$ that is played from the next period on if i rejects the proposal in the first period. From the argument in the previous paragraph, we thus have $U_j^a \geq (1 - \delta)\gamma\Delta r_j + \delta V_j(\sigma^r)$ for all $j \in W$. Similarly, every proposal that may successfully be made by the penultimate proposer in the first period under σ^m (both on and off the path) must also be accepted by some member i of W . Her payoff (and therefore the payoff of all members of W) from accepting must be at least as large as the payoff from rejecting which, as previously shown, must be at least $(1 - \delta)\gamma\Delta r_i + \delta \inf \{V_i(\sigma) : \sigma \in \mathcal{E}(R, 1)\}$. Applying the same argument recursively, we obtain that the acceptance of any proposal in the first period must give a payoff of at least $(1 - \delta)\gamma\Delta r_i + \delta \inf \{V_i(\sigma) : \sigma \in \mathcal{E}(R, 1)\}$ for all $i \in W$. Hence,

$$V_i(\sigma^m) \geq (1 - \delta)\gamma\Delta r_i + \delta \inf \{V_i(\sigma) : \sigma \in \mathcal{E}(R, 1)\},$$

for all $i \in W$. Taking the limit as $m \rightarrow \infty$ and recalling the definition of $\{\sigma^m\}$, we obtain $\gamma\Delta r_i = \inf \{V_i(\sigma) : \sigma \in \mathcal{E}(R, 1)\}$ (since $\gamma\Delta r_i$ is maximum feasible payoff for a player $i \in W$ when R is good with probability one). This in turn implies that R must be implemented with probability one in every period of every equilibrium of $\Gamma(R, 1)$.

The argument for the case where $L \in \mathcal{D}$ is analogous.

We now turn to part (ii). To prove the second part of the lemma, we proceed in three steps: first, we show that the infimum of every player i 's equilibrium payoff in $\Gamma(S, \alpha_k)$ converges to Δs_i as $k \rightarrow \infty$; then, we show that for sufficiently large k , alternative S is implemented in every period of $\Gamma(S, \alpha_k)$; finally, we use the previous result to complete the proof of the lemma.

Let $\alpha_k \in A \setminus \{1\}$; and let $\mathcal{E}(S, \alpha_k)$ be the set of PBEs of $\Gamma(S, \alpha_k)$. Every period of $\Gamma(S, \alpha_k)$ begins with a belief α that alternative R is good; then, either S is implemented, in which case committee member i receives a payoff of Δs_i ; or R is implemented, in which case i 's expected payoff is $\alpha\gamma\Delta r_i$. Therefore, every strategy profile σ yields an expected payoff of the form

$$V_i^k(\sigma) \equiv \Delta \left[\beta_s^k(\sigma) s_i + \beta_1^k(\sigma) \gamma r_i + \sum_{\ell=k}^{\infty} \beta_\ell^k(\sigma) \alpha_\ell \gamma r_i \right]$$

to player i in $\Gamma(S, \alpha_k)$, where $\beta_s^k(\sigma) + \beta_1^k(\sigma) + \sum_{\ell=k}^{\infty} \beta_\ell^k(\sigma) = 1$ and $\beta_s^k(\sigma), \beta_1^k(\sigma), \beta_\ell^k(\sigma) \in [0, 1]$ for each $\ell = k, k+1, \dots$. Moreover, as R must have been successfully tried at least once to be known to be good, $\beta_1^k(\cdot)$ is bounded above by $\alpha_k \gamma$. Coupled with the fact that $\alpha_\ell \leq \alpha_k$ for all $\ell \geq k$, this implies that $\lim_{k \rightarrow \infty} M_i^k \equiv \sup_{\sigma \in \mathcal{E}(S, \alpha_k)} \left[\beta_1^k(\sigma) + \sum_{\ell=k}^{\infty} \beta_\ell^k(\sigma) \alpha_\ell \right] \gamma r_i$ for each $i \in N$. This in turn implies that there is a null sequence $\{\varepsilon^k\}$ in $\mathbb{R}_+$ such that, for all $\sigma \in \mathcal{E}(S, \alpha_k)$, we have

$$\max_{i \in N} |V_i^k(\sigma) - \beta_s^k(\sigma) s_i \Delta| \leq \Delta \max_{i \in N} M_i^k < \varepsilon^k,$$

for every $k \in \mathbb{N}$. Now for each $k \in \mathbb{N}$, let $\{\sigma^{k,m}\}$ be a sequence in $\mathcal{E}(S, \alpha_k)$ such that $\lim_{m \rightarrow \infty} \beta_s^k(\sigma^{k,m}) = \inf_{\sigma \in \mathcal{E}(S, \alpha_k)} \beta_s^k(\sigma)$. As $\sigma^{k,m}$ is an equilibrium of $\Gamma(S, \alpha_k)$, there must be at least one committee member, say i_k , such that

$$V_{i_k}^k(\sigma^{k,m}) \geq (1 - \delta) \Delta s_{i_k} + \delta \inf_{\sigma \in \mathcal{E}(S, \alpha_k)} V_{i_k}^k(\sigma);$$

otherwise some player would have a profitable deviation in the first period of $\Gamma(S, \alpha_k)$. It follows that

$$\beta_s^k(\sigma^{k,m}) \Delta s_{i_k} + \varepsilon^k \geq (1 - \delta) \Delta s_{i_k} + \delta \left[\inf_{\sigma \in \mathcal{E}(S, \alpha_k)} \beta_s^k(\sigma) \Delta s_{i_k} - \varepsilon^k \right].$$

Taking the limit as $m \rightarrow \infty$, we obtain $\inf_{\sigma \in \mathcal{E}(S, \alpha_k)} \beta_s^k(\sigma) \geq 1 - 2(\varepsilon^k / \Delta s_{i_k})$. This implies that $\inf_{\sigma \in \mathcal{E}(S, \alpha_k)} \beta_s^k(\sigma)$ converges to one as $k \rightarrow \infty$ and, therefore, that there exists a null sequence $\{\eta^k\}$ such that $\lim_{k \rightarrow \infty} \max_{i \in N} \left| \inf_{\sigma \in \mathcal{E}(S, \alpha_k)} V_i^k(\sigma) - \Delta s_i \right| < \eta^k$, for all $k \in \mathbb{N}$, thus completing the first step of the argument.

We now turn to the second step of the proof. Observe first that, as $(1 - \delta)(s_i - \alpha_k \gamma r_i) \Delta - \delta \eta^k$ converges to $(1 - \delta) \Delta s_i > 0$ as $k \rightarrow \infty$, there is a sufficiently large $K \in \mathbb{N}$ such that $(1 - \delta)(s_i - \alpha_k \gamma r_i) \Delta - \delta \eta^k > 0$, for all $k \geq K$. Let $k \geq K$, and suppose that $\Gamma(S, \alpha_k)$ has an equilibrium in which alternative R is implemented with positive probability. Consider the first period of $\Gamma(S, \alpha_k)$ in which R may be implemented. Every decisive voter i 's benefit from rejecting any proposal to change S to R is bounded below by $(1 - \delta)(s_i - \alpha_k \gamma r_i) \Delta + \delta [(\Delta s_i - \eta^k) - \Delta s_i] > 0$, where the bracketed term represents the difference between the lower and upper bounds on i 's continuation payoffs from rejecting R and accepting it, respectively. (Recall that each committee member i 's maximum payoff is

Δs_i when the belief is smaller than or equal to $\hat{\alpha}_n$.) Hence, then every proposal to amend S to R is rejected in any equilibrium of $\Gamma(S, \alpha_k)$. We thus have $V_i^k(\sigma) = \Delta s_i$, for all $i \in N$ and all $\sigma \in \mathcal{E}(S, \alpha_k)$.

If $\alpha_K > \hat{\alpha}_n$, then Lemma 1(ii) follows immediately from the previous paragraph; so suppose that $\alpha_K \leq \hat{\alpha}_n$. To complete the proof of the result, consider the first period of $\Gamma(S, \alpha_K)$. If alternative R is implemented, then the expected payoff to each committee member i is $[1 - \delta(1 - \gamma\Delta)]\alpha_K\gamma\Delta r_i + \delta(1 - \alpha_K\gamma\Delta)\Delta s_i < \Delta s_i$, where the inequality follows from $\alpha_K \leq \hat{\alpha}_n$ and the definition of the committee members' optimal cutoffs in Subsection 3.2; if alternative S is instead implemented, then her expected payoff will be a convex combination of $[1 - \delta(1 - \gamma\Delta)]\alpha_K\gamma\Delta r_i + \delta(1 - \alpha_K\gamma\Delta)\Delta s_i$ (if R is implemented with positive probability in a future period) and Δs_i , with a positive coefficient on the latter. Therefore, all committee members are strictly better off implementing R : they all reject proposals to amend S to R (when decisive). Hence, $V_i^k(\sigma) = \Delta s_i$, for all $i \in N$, $k \geq K$ and $\sigma \in \mathcal{E}(S, \alpha_k)$. Applying the same argument recursively from belief α_{K-1} to belief $\hat{\alpha}_n$ we obtain that, for all $\alpha_k \leq \hat{\alpha}_n$, $\Gamma(S, \alpha_k)$ has a unique equilibrium outcome: Alternative S is implemented in every period. By the same logic, the same is also true in game $\Gamma(S, \alpha_k)$, $\alpha_k \leq \hat{\alpha}_n$. In such a game, every decisive voter receives her largest possible payoff Δs_i if she accepts a proposal to change the status quo R to S , since the latter will then never be amended. Any such a proposal must therefore be successful and, as S is the ideal policy of all players, some proposer must successfully propose it in equilibrium.

Finally, S being the ideal alternative of all the players, it is easy to construct an equilibrium in which all players always propose alternative S (conditional on being recognized to propose), accept any proposal to change status quo R to alternative S , and reject any proposal to amend status quo S .

B Proof of Proposition 2: Collegial Rules

Section B in the main text contains a proof of Proposition 2 for cases where the voting rule is noncollegial or unanimity. This section covers the other cases.

(i) **Unanimity rule.** Suppose now that $\mathcal{D} = \{N\}$. As in the proof of Proposition 1, we denote by $\Gamma(p \mid \alpha)$ the continuation game that begins with status quo $p \in \{R, S\} \times [0, 1] \times X$ and belief $\alpha \in A$.

Lemma B1. *Suppose $\mathcal{D}$ is unanimity rule and $\hat{\tau} = 1$. Then there exists $\bar{\Delta}_1 > 0$ such that, for all $\Delta < \bar{\Delta}_1$ and all $(\tau, x) \in [0, 1] \times X$:*

- (i) $\Gamma(R, \tau, x \mid 1)$ has a renegotiation-proof equilibrium and, in any such equilibrium, each committee member i 's payoff is $w_i(R, \tau, x \mid 1) \equiv \gamma \Delta [(1 - \tau)r_i + \tau x_i \bar{r}]$; and
- (ii) the set of renegotiation-proof equilibrium payoffs for $\Gamma(S, \tau, x \mid 1)$ is the simplex $\{(w_1, \dots, w_n) \in \mathbb{R}^n : \sum_{i=1}^n w_i = V^*(1) \text{ and } w_i \geq [(1 - \tau)s_i + \tau x_i \bar{s}] \Delta, \forall i \in N\}$.

Proof. Let $\bar{\Delta}_1 \equiv \sup \{\Delta \in \mathbb{R}_+ : (1 - e^{-\rho \Delta}) < e^{-\rho \Delta} (n - 2)\} > 0$. Henceforth, we assume that $\Delta < \bar{\Delta}_1$.

Let $(\tau, x) \in [0, 1] \times X$, let $w^0 \in \{(w_1, \dots, w_n) \in \mathbb{R}^n : \sum_{i=1}^n w_i = V^*(1) \text{ and } w_i \geq [(1 - \tau)s_i + \tau x_i \bar{s}] \Delta\}$. To prove the lemma, we will construct a renegotiation-proof equilibrium σ for $\Gamma(S, \tau, x \mid 1)$ in which every committee member i receives w_i^0 , thus establishing part (ii). In that equilibrium, the optimal stopping rule will be implemented in every period both on and off the path. As $\Gamma(R, \tau, x \mid 1)$ is itself a continuation game in $\Gamma(S, \tau, x \mid 1)$, this will also establish that $\Gamma(R, \tau, x \mid 1)$ has a renegotiation-proof equilibrium. If the status quo is (R, τ, x) and the belief is equal to 1, then each committee member i can obtain a payoff of $w_i(R, \tau, x \mid 1)$ by rejecting any future proposal to amend the status quo. As the payoff vector $(w_j(R, \tau, x \mid 1))_{j \in N}$ is in the Pareto frontier (and $\mathcal{D}$ is unanimity rule), it follows that $w_i(R, \tau, x \mid 1)$ is i 's payoff in any (renegotiation-proof) equilibrium for $\Gamma(R, \tau, x \mid 1)$.

We begin with an intuitive description of the equilibrium σ . Alternative R is implemented in each period (both on and off the path). As the belief is equal to 1, this implies that payoff vectors are Pareto optimal in every continuation game. Once S has been implemented, all proposers pass in all future periods, irrespective of the tax rate and distribution of revenues. If R has not yet been implemented, then behavior is determined by a set of n “phases,” each corresponding to one committee member in N . In committee member i 's phase, every proposer successfully offers a policy that gives a payoff of $\gamma \Delta \bar{r} - \sum_{j \neq i} [(1 - \tau')s_j + \tau y_j \bar{s}] \Delta$ to i and a payoff of $[(1 - \tau')s_j + \tau y_j \bar{s}] \Delta$ to each commit-

tee member $j \neq i$, where (S, τ', y) is the status quo policy. The idea is that i receives her “reward payoff” and the others their “punishment payoffs.” If a proposer, say i , deviates, then every committee member (other than i) rejects her proposal and the game transitions to the phase of the first committee member who rejected the proposal. If voter i rejects a proposal which she should have accepted, then the game moves to the another committee member’s phase.

We now turn to the formal definition of σ for the continuation game $\Gamma(S, \tau, x \mid 1)$. As in the case of noncollegial rules, we divide each period into n “parts,” each consisting of a proposal stage and the n voting stages that follow it. Changes of phases can only occur at the end of these parts. A phase is formally represented by a pair $(\ell, i) \in \{1, \dots, n\} \times \{0, 1, \dots, n\}$. In every period that begins with status quo $p = (a, \tau', y)$ and an order of proposers $(\pi_1, \dots, \pi_n)$, σ prescribes the following behavior in phase (ℓ, i) :

(a) If $a = R$, then proposer π_ℓ passes; and if $a = S$, then she offers policy $(R, 1, y^i)$, where

$$y_j^i \equiv \begin{cases} w_j^0/V^*(1) & \text{if } i = 0, \\ w_j(p \mid 1)\Delta/V^*(1) & \text{if } i \neq 0 \text{ \& } j \neq i, \\ [V^*(1) - \sum_{k \neq i} w_k(p \mid 1)]/V^*(1) & \text{if } i \neq 0 \text{ \& } j = i, \end{cases}$$

for all $j \in N$;

(b) if $a = S$ and π_ℓ offered $(R, 1, y^i)$, then every voter accepts it (irrespective of the previous voters’ actions);

(c) if $a = R$ and π_ℓ proposed some policy $p' = (a', \tau'', z) \neq p$, then voter $j \in N$ votes to accept it if and only if

$$w_j(p \mid 1) < (1 - \delta)w_j(p' \mid 1) + \delta \begin{cases} w_j(p' \mid 1) & \text{if } a' = R, \\ y_j^i V^*(1) & \text{if } a' = S. \end{cases};$$

(d) if $a = S$ and π_ℓ proposed some policy $p' \notin \{(R, 1, y^i), p\}$, then every voter j acts according to the following rules:

(d1) if any previous voter has already rejected p' , then j also rejects it;

(d2) if she is the n th voter and p' has not yet been rejected by any voter, then she

accepts p' if and only if the following holds

$$\begin{cases} (1 - \delta)w_j(p' | 1) + \delta y_j^{\pi_\ell} V^*(1) > (1 - \hat{\delta}_\ell)w_j(p | 1) + \hat{\delta}_\ell y_j^j V^*(1) & \text{if } j \neq \pi_\ell, \\ (1 - \delta)w_j(p' | 1) + \delta y_j^{\pi_\ell} V^*(1) > (1 - \hat{\delta}_\ell)w_j(p | 1) + \hat{\delta}_\ell y_j^{\hat{i}} V^*(1) & \text{if } j = \pi_\ell, \end{cases}$$

where $\hat{i} \equiv \min N \setminus \{\pi_\ell\}$ and

$$\hat{\delta}_\ell \equiv \begin{cases} 1 & \text{if } \ell \in \{1, \dots, n-1\}, \\ \delta & \text{if } \ell = n; \end{cases}$$

(d3) if she is the k th voter, $k < n$, and p' has not yet been rejected by any voter, then she accepts p' if and only if all the remaining voters will also accept it and the following holds

$$\begin{cases} (1 - \delta)w_j(p' | 1) + \delta y_j^{\pi_\ell} V^*(1) > (1 - \hat{\delta}_\ell)w_j(p | 1) + \hat{\delta}_\ell y_j^j V^*(1) & \text{if } j \neq \pi_\ell, \\ (1 - \delta)w_j(p' | 1) + \delta y_j^{\pi_\ell} V^*(1) > (1 - \hat{\delta}_\ell)w_j(p | 1) + \hat{\delta}_\ell y_j^{\hat{i}} V^*(1) & \text{if } j = \pi_\ell. \end{cases}$$

Observe that, from (a) and (b) above, every committee member j 's continuation value at the start of phase (ℓ, i) is $w_j(p | 1)$ if $a = R$, and $y_j^i V^*(1)$ if $a = S$.

In period 1, play begins in phase $(1, 0)$. Then in every period, at the end of any part that began with status quo $p = (a, \tau', y)$ and in some phase $(\ell, i) \in \{1, \dots, n\} \times \{0, 1, \dots, n\}$:

(t1) If $a = R$ and π_ℓ passed, then the game transitions to phase $(\ell + 1, i)$ (we set $\ell + 1 = 1$ whenever $\ell = n$);

(t2) if $a = S$ and π_ℓ proposed $(R, 1, y^i)$ which was accepted, then the game transitions to phase $(1, i)$;

(t3) if $a = S$ and π_ℓ proposed $(R, 1, y^i)$ which was rejected, then the game transitions to phase $(\ell + 1, j + 1)$ (set $j + 1 = 1$ whenever $j = n$), where j is the first voter who rejected it;

(t4) if $a = R$ and π_ℓ proposed some policy $p' \neq p$, then the game transitions to phase $(1, i)$ if p' was accepted, and to phase $(\ell + 1, i)$ otherwise;

(t5) if $a = S$ and π_ℓ proposed some policy $p' \notin \{(R, 1, y^i), p\}$ which was accepted, then the game transitions to phase $(\ell + 1, \pi_\ell)$;

(t6) if $a = S$ and π_ℓ proposed some policy $p' \notin \{(R, 1, y^i), p\}$ which was rejected by some voter in $N \setminus \{\pi_\ell\}$, then the game transitions to phase $(\ell + 1, j)$, where j is the first voter in $N \setminus \{\pi_\ell\}$ who rejected p' ; and

(t7) if $a = S$ and π_ℓ proposed some policy $p' \notin \{(R, 1, y^i), p\}$ which was only rejected by π_ℓ herself, then the game transitions to phase $(\ell + 1, \hat{i})$;

(t8) if $a = S$ and π_ℓ passed, then the game transitions to phase $(\ell + 1, \pi_\ell + 1)$.

We now verify that for $\Delta < \bar{\Delta}$, σ is a PBE. Take an arbitrary committee member $j \in N$, and consider a voting stage with status quo $p = (a, \tau', y)$ in some phase (ℓ, i) . Suppose first that $a = S$ and π_ℓ proposed $(R, 1, y^i)$, so that σ prescribes j to accept this proposal (see (b)). If some previous voter has already rejected the proposal, then j has trivially no profitable deviation: her decision will have no impact on her payoff. If she is the first voter, or if all previous voters have accepted the proposal, then her decision does impact her payoff. If she accepts $(R, 1, y^i)$ then, from (b) and (t2) above, she receives a payoff of $y_j^i V^*(1)$; if she rejects $(R, 1, y^i)$ then, from (t3), the game transitions to phase $(\ell + 1, j + 1)$ and she receives $(1 - \hat{\delta})w_j(p | 1) + \hat{\delta}y_j^{j+1}V^*(1)$. Since

$$y_j^i V^*(1) \geq w_j(p | 1) = (1 - \hat{\delta}_\ell)w_j(p | 1) + \hat{\delta}_\ell y_j^{j+1} V^*(1) ,$$

she is better off accepting.

Suppose now that $a = R$ and π_ℓ proposed some policy $p' \neq p$. If p' is accepted then, from (t4), the game transitions to phase $(1, i)$ and committee member j receives $(1 - \delta)w_j(p' | 1) + \delta \begin{cases} w_j(p' | 1) & \text{if } a' = R , \\ y_j^i V^*(1) & \text{if } a' = S; \end{cases}$ if p' is rejected then she receives $w_j(p | 1)$. It follows from (c) that, under σ , her decision is optimal whenever she is decisive. Hence, she cannot profitably deviate from σ .

Finally, suppose that $a = S$ and π_ℓ proposed some policy $p' \notin \{(R, 1, y^i), p\}$:

- If some voter in $N \setminus \{\pi_\ell, j\}$ has already rejected the proposal then, from (t6), her decision will not have any impact on her payoff and is therefore optimal.
- If π_ℓ is the only voter who has already rejected p' , then the choice of voter $j \neq \pi_\ell$ has no impact on her stage-game payoff in this period but will impact the transition to the next phase. It follows from (t6) and (t7) that she cannot improve on rejecting, which is the action prescribed by σ (see (d1)).

- If p' has not yet been rejected and j is the n th voter, then it follows from (d2) and (t5)-(t7) that σ prescribes her to accept p' if and only if she is strictly better off doing so.

The same is true if j is not the last voter and she anticipates that all the remaining voters will accept p' — see (d3).

- If p' has not yet been rejected, j is not the last voter and she anticipates the some of the remaining voters will reject p' , then her choice has no impact on the policy that will be implemented in the current period. If, in addition, $j = \pi_\ell$ then her choice does not have any impact on her continuation value either and, therefore, rejecting is optimal. If instead $j \neq \pi_\ell$, then her decision will impact the transition to the next phase. It follows from (t6) and (t7) that she is better off rejecting, as prescribed by σ . (She can only be indifferent if π_ℓ is the only other voter who will reject p' , and $j = \hat{i}$.)

This proves that deviations in voting stages are unprofitable. We now turn to proposal stages. Consider the proposal stage of any phase (ℓ, i) that begins with a status quo $p = (a, \tau', y)$. Suppose first that $a = R$. If $j = \pi_\ell$ passes, as prescribed by σ , then from (t1) she receives a payoff of $y_j^i V^*(1)$. If she deviates by proposing a policy $p' \neq p$ then, from (c), her proposal will be rejected: as the payoff vector $(w_k(p | 1))_{k \in N}$ belongs to the Pareto frontier, it is impossible to offer every committee member k a higher payoff than $w_k(p | 1)$. It then follows from (a) that proposer j gets a payoff of $w_j(p | 1) \leq y_j^i V^*(1)$ (with a strict inequality if and only if $i = j$). Hence, j cannot profitably deviate from passing.

Suppose now that $a = S$. If $j = \pi_\ell$ proposes $(R, 1, y^i)$, as prescribed by σ , then from (b), (t2) and (a), she receives a payoff of $y_j^i V^*(1)$. If she deviates by passing, then she receives $(1 - \delta)w_j(p | 1)$ in the current period and the game then transitions to phase $(\ell + 1, j + 1)$ (see (t8)). She thus receives

$$(1 - \hat{\delta}_\ell)w_j(p | 1) + \hat{\delta}_\ell y_j^{j+1} V^*(1) = w_j(p | 1) \leq y_j^i V^*(1)$$

(with a strict inequality if and only if $i = j$). Hence, passing is not a profitable deviation. Finally, if she deviates by proposing a policy $p' \neq (R, 1, y^i)$ then, from (c), her proposal will be rejected. To see this, observe that from (d), she would have to offer more than $(1 - \hat{\delta}_\ell)w_k(p | 1) + \hat{\delta}_\ell y_k^k V^*(1)$ to all the other committee members $k \neq j$, and more than $w_j(p | 1)$ to herself. Summing across the committee members and rearranging terms, a successful proposal would have to generate a total sum of payoffs that exceeds $(1 - \hat{\delta}_\ell)\bar{s}\Delta +$

$\hat{\delta}_\ell[V^*(1) + (n-2)[V^*(1) - \bar{s}\Delta]] = (1 - \hat{\delta}_\ell)\bar{s}\Delta + \hat{\delta}_\ell[\gamma\bar{r} + (n-2)[\gamma\bar{r} - \bar{s}]]\Delta$. As aggregate payoffs are bounded above by $V^*(1) = \gamma\Delta\bar{r}$, this would require that

$$(1 - \hat{\delta}_\ell)\bar{s} + \hat{\delta}_\ell[\gamma\bar{r} + (n-2)[\gamma\bar{r} - \bar{s}]] < \gamma\bar{r} \quad (\text{B2})$$

or, equivalently, $1 - \hat{\delta}_\ell < \hat{\delta}_\ell(n-2)$. This is impossible since $\Delta < \bar{\Delta}_1$. Moreover, it follows from (t6)-(t7) that the game will then transition to phase $(\ell+1, j)$. As a result, j obtains a payoff of

$$(1 - \hat{\delta}_\ell)w_j(p \mid 1) + \hat{\delta}_\ell y_j^k V^*(1) = w_j(p \mid 1) \leq y_j^i V^*(1)$$

(with $k \neq j$) if she deviates by making a proposal $p' \neq (R, 1, y^i)$. Hence, such a deviation is unprofitable. $\square$

Lemma B2. *Suppose $\mathcal{D}$ is unanimity rule and $\hat{\tau} = 1$. Then there exists $\bar{\Delta}_2 > 0$ such that, for all $\Delta < \bar{\Delta}_2$, all beliefs $\alpha_k \leq \alpha^*$ and all $(\tau, x) \in [0, 1] \times X$:*

- (i) $\Gamma(S, \tau, x \mid \alpha_k)$ has a renegotiation-proof equilibrium and, in any such equilibrium, each committee member i 's payoff is $w_i(S, \tau, x \mid \alpha_k) \equiv \alpha_k \gamma \Delta [(1 - \tau)s_i + \tau x_i \bar{s}]$; and
- (ii) the set of renegotiation-proof equilibrium payoffs for $\Gamma(R, \tau, x \mid \alpha_k)$ is the simplex $\{(w_1, \dots, w_n) \in \mathbb{R}^n : \sum_{i=1}^n w_i = \bar{s}\Delta \text{ and } w_i \geq \alpha_k \gamma \Delta [(1 - \tau)r_i + \tau x_i \bar{r}], \forall i \in N\}$.

Proof. An application of l'Hôpital's rule gives

$$\lim_{\Delta \rightarrow 0} \alpha^* = \frac{\rho \bar{s}}{\gamma[(\rho + \gamma)\bar{r} - \bar{s}]} ,$$

so that

$$\bar{s} - \gamma\bar{r} \lim_{\Delta \rightarrow 0} \alpha^* = \frac{\bar{s}(\gamma\bar{r} - \bar{s})}{(\rho + \gamma)\bar{r} - \bar{s}} > 0 .$$

It follows that

$$\tilde{\Delta}_2 \equiv \sup \{ \Delta > 0 : [1 - \delta(1 - \alpha^* \gamma \Delta)] \bar{s} - [1 - \delta(1 - \gamma \Delta)] \alpha^* \gamma \bar{r} < \delta(1 - \alpha^* \gamma \Delta)(n-2)[\bar{s} - \alpha^* \gamma \bar{r}] \}$$

is well-defined and positive. Observe that if $\Delta < \tilde{\Delta}_2$, then the following inequality holds for all beliefs $\alpha_k \leq \alpha^*$:

$$[1 - \delta(1 - \alpha_k \gamma \Delta)] \bar{s} - [1 - \delta(1 - \gamma \Delta)] \alpha_k \gamma \bar{r} < \delta(1 - \alpha_k \gamma \Delta)(n-2)[\bar{s} - \alpha_k \gamma \bar{r}] .$$

Henceforth, we assume that $\Delta < \bar{\Delta}_2 \equiv \min\{\bar{\Delta}_1, \tilde{\Delta}_2\}$.

To prove the lemma, one can use an equilibrium construction that parallels that in the proof of Lemma B1. In this equilibrium, when the status quo is of the form $p = (a, \tau', y)$ and the belief is $\alpha \leq \alpha_k$, committee member i 's reward payoff is $V^*(\alpha) - \sum_{j \neq i} w_j(p \mid \alpha) = \bar{s}\Delta - \sum_{j \neq i} \alpha\gamma\Delta[(1 - \tau')r_j + \tau'y_j\bar{s}]$ and her punishment payoff is $w_i(p \mid \alpha) = \alpha\gamma\Delta[(1 - \tau')r_i + \tau'y_i\bar{s}]$. If the belief becomes equal to one, then the equilibrium described in Lemma B1 is played. The argument is then exactly the same. In particular, the key condition (B2), necessary for proposers to have profitable deviations, now becomes

$$[1 - \hat{\delta}_\ell(1 - \gamma\Delta)]\alpha\gamma\Delta\bar{r} + \hat{\delta}_\ell(1 - \alpha\gamma\Delta)[V^*(\alpha) + (n - 2)(V^*(\alpha) - \alpha\gamma\Delta\bar{r})] < V^*(\alpha)$$

or, equivalently,

$$\hat{\delta}_\ell(1 - \alpha_k\gamma\Delta)(n - 2)[\bar{s} - \alpha\gamma\bar{r}] < [1 - \hat{\delta}_\ell(1 - \alpha\gamma\Delta)]\bar{s} - [1 - \hat{\delta}_\ell(1 - \gamma\Delta)]\alpha\gamma\bar{r}.$$

This cannot hold since $\Delta < \bar{\Delta}_2$. □

Lemma B3. *Suppose $\mathcal{D}$ is unanimity rule and $\hat{\tau} = 1$. Then there exists $\bar{\Delta}_2 > 0$ such that, for all $\Delta < \bar{\Delta}_2$, all beliefs $\alpha_k \in [\alpha_1, \alpha^*)$ and all $(\tau, x) \in [0, 1] \times X$:*

- (i) $\Gamma(a, \tau, x \mid \alpha_k)$, $a \in \{R, S\}$, has a renegotiation-proof equilibrium that sustains the optimal stopping rule; and
- (ii) $\Gamma(R, 1, x \mid \alpha_k)$ has a renegotiation-proof equilibrium in which each committee member i 's expected payoff is $x_i V^*(\alpha_k)$.

Proof. Let $\bar{\Delta}_2$ be defined as in Lemma B2. To prove Lemma B3, consider first the simple variant on the standard Baron-Ferejohn model, denoted $\mathbf{G}(a, \tau, x \mid \alpha_k)$, in which the policy space is not the unit simplex but $\{(w_1, \dots, w_n) \in \mathbb{R}^n : \sum_{i=1}^n w_i \leq V^*(\alpha_k) \text{ and } w_i \geq w_i(S, \tau, x \mid \alpha_k), \forall i \in N\}$, and the probability that committee member i is selected to propose is equal to the probability q_i that she is the last proposer in our model. It is well known that this game has a (pure strategy) stationary subgame perfect equilibrium, in which the selected proposer makes the same (successful) proposal, $w^i(S, \tau, x \mid \alpha_k)$, in every period.

Let $k^* \in \mathbb{N}$ be implicitly defined by $\alpha_{k^*} = \alpha^*$, and let $(\tau, x) \in [0, 1] \times X$. Consider a strategy profile σ^{k^*-1} for $\Gamma(R, \tau, x \mid \alpha_{k^*-1})$ that prescribes the following behavior in any period that begins with a status quo $p = (a, \tau', y)$, a belief $\alpha \in \{\alpha_k \in A: k \geq k^* - 1\} \cup \{1\}$, and an order of proposers $(\pi_1, \dots, \pi_n)$:

a) If $\alpha = \alpha_{k^*-1}$ and $a = R$, then all proposers pass (irrespective of the previous history of play);

b) if $\alpha = \alpha_{k^*-1}$, $a = R$, and some proposer has offered a policy $p' = (a', \tau'', z) \neq p$, then voter i votes to accept p' if and only if

$$\frac{w_i(p \mid \alpha_{k^*-1})}{\alpha_{k^*-1} \gamma \Delta \bar{r}} V^*(\alpha_{k^*-1}) < \begin{cases} \frac{w_i(p' \mid \alpha_{k^*-1})}{\alpha_{k^*-1} \gamma \Delta \bar{r}} V^*(\alpha_{k^*-1}) & \text{if } a' = R, \\ (1 - \delta) s_i \Delta + \delta V^*(\alpha_{k^*-1}) \sum_{j=1}^n q_j w_i^j(p' \mid \alpha_{k^*-1}) & \text{if } a' = S; \end{cases}$$

c) if $\alpha = \alpha_{k^*-1}$ and $a = S$, then each proposer π_ℓ , $\ell < n$, passes and proposer π_n offers policy $(R, 1, y^{\pi_n}(S, \tau', y \mid \alpha_{k^*-1}))$;

d) if $\alpha = \alpha_{k^*-1}$, $a = S$, and some proposer has offered a policy $p' = (a', \tau'', z) \neq p$, then voter i makes the same decision as when she is offered the following policy in the stationary subgame perfect equilibrium of $\mathbf{G}(S, \tau', y \mid \alpha_{k^*-1})$:

$$\begin{cases} \frac{w_i(p' \mid \alpha_{k^*-1})}{\alpha_{k^*-1} \gamma \Delta \bar{r}} V^*(\alpha_{k^*-1}) & \text{if } a' = R, \\ (1 - \delta) s_i \Delta + \delta V^*(\alpha_{k^*-1}) \sum_{j=1}^n q_j w_i^j(p' \mid \alpha_{k^*-1}) & \text{if } a' = S; \end{cases}$$

e) if $\alpha \leq \alpha^*$ and $a = R$, then the committee plays an equilibrium of $\Gamma(R, \tau', y \mid \alpha)$ in which each committee member i 's payoff is $\frac{w_i(p \mid \alpha)}{\alpha \gamma \Delta \bar{r}} V^*(\alpha)$ (see Lemma B2(ii)); if $\alpha \leq \alpha^*$ and $a = S$, then the committee plays an equilibrium of $\Gamma(R, \tau', y \mid \alpha)$ in which each committee member i 's payoff is $w_i(p \mid \alpha)$ (see Lemma B2(i));

d) if $\alpha = 1$ and $a = R$, then the committee plays an equilibrium of $\Gamma(R, \tau', y \mid 1)$ in which each committee member i 's payoff is $w_i(p \mid 1)$ (see Lemma B1(i)); if $\alpha = 1$ and $a = S$, then the committee plays an equilibrium of $\Gamma(S, \tau', y \mid 1)$ in which each committee member i 's payoff is $\frac{w_i(p \mid 1)}{\bar{s} \Delta}$ (see Lemma B1(ii)).

It is readily checked that σ^{k^*-1} is a PBE for $\Gamma(R, \tau, x \mid \alpha_{k^*-1})$. In particular, the acceptance condition in case b) compares the voter's continuation value from rejecting the proposal (left side of the inequality) with her continuation value from accepting it (right

side). As the optimal stopping rule is implemented in case of rejection (and the voting rule is unanimity), at least one voter must vote to reject the proposal. It follows that any proposal is unsuccessful in case a) and, therefore, passing is optimal for all proposers. In the other cases, any deviation is by construction unprofitable. Moreover, as the optimal stopping rule is implemented in every continuation game both on and off the equilibrium path, σ^{k^*-1} is a renegotiation-proof equilibrium. As $\Gamma(S, \tau, x \mid \alpha_{k^*-1})$ is a continuation game of $\Gamma(R, \tau, x \mid \alpha_{k^*-1})$, the restriction of σ^{k^*-1} to $\Gamma(S, \tau, x \mid \alpha_{k^*-1})$ is also a renegotiation-proof equilibrium. To complete the proof of the lemma for the case where $k = k^* - 1$, observe that if $\tau = 1$, then each committee member i 's payoff in equilibrium σ^{k^*-1} is $x_i V^*(\alpha_{k^*-1})$.

To obtain the result for any $k \in \{1, \dots, k^*-1\}$, one can then proceed recursively: having obtained an equilibrium σ^{k+1} for every continuation game of the form $\Gamma(a, \tau', y \mid \alpha_{k+1})$, one can apply the same construction as above at belief α_k to obtain a renegotiation-proof equilibrium σ^k for every game $\Gamma(a, \tau, x \mid \alpha_k)$. $\square$

Lemma B4. *Suppose $\mathcal{D}$ is unanimity rule and $\hat{\tau} = 1$. Then there exists $\bar{\Delta} > 0$ such that, for all $\Delta < \bar{\Delta}$, the set of renegotiation-proof equilibrium payoffs for $\Gamma(S, 0, x^0 \mid \alpha_0)$ is the simplex $\{(w_1, \dots, w_n) \in \mathbb{R}^n : \sum_{i=1}^n w_i = V^*(\alpha_0) \text{ and } w_i \geq s_i \Delta, \forall i \in N\}$.*

Proof. It follows from Assumption A1 that, for sufficiently small Δ , $V^*(\alpha_0)/\Delta > \bar{s}$. Therefore, the threshold

$$\tilde{\Delta}_3 \equiv \{\Delta > 0 : \delta(n-2)[(V^*(\alpha_0)/\Delta) - \bar{s}] < (1-\delta)[(V^*(\alpha_0)/\Delta) - \bar{s}]\}$$

is well-defined and positive. Henceforth, we assume that $\Delta < \bar{\Delta} \equiv \min\{\bar{\Delta}_2, \tilde{\Delta}_3\}$.

To prove the lemma, one can use again an equilibrium construction that parallels that in the proof of Lemma B1. In this equilibrium, when the status quo is of the form $p = (S, \tau', y)$ and the belief is α_0 , committee member i 's reward payoff is $V^*(\alpha_0) - \sum_{j \neq i} w_j(p \mid \alpha_0) = \bar{s}\Delta - \sum_{j \neq i} s_j \Delta$ and her punishment payoff is $w_i(p \mid \alpha_0) = s_i \Delta$. If the belief becomes equal to α_1 , then an equilibrium as described in Lemma B3 is played — in particular, the equilibrium described in Lemma B3(ii) if the status quo is of the form $(R, 1, x)$ for some $x \in X$ — and if the belief becomes equal to one, then the equilibrium described in Lemma

B1 is played. The argument is then exactly the same. In particular, the key condition (B2), necessary for proposers to have profitable deviations, is now

$$(1 - \hat{\delta}_\ell)\bar{s}\Delta + \hat{\delta}_\ell[V^*(\alpha_0) + (n-2)(V^*(\alpha_0) - \bar{s}\Delta)] < V^*(\alpha_0)$$

or, equivalently,

$$\hat{\delta}_\ell(n-2) \left[\frac{V^*(\alpha_0)}{\Delta} - \bar{s} \right] < (1 - \hat{\delta}_\ell) \left[\frac{V^*(\alpha_0)}{\Delta} - \bar{s} \right].$$

As $\Delta < \bar{\Delta}$, this inequality cannot hold. $\square$

Let $\Delta < \bar{\Delta}$, where $\bar{\Delta} > 0$ is the threshold defined in Lemma B4. To complete the proof for the unanimity case, observe that in any PBE, each committee member i 's expected payoff must be greater than or equal to $s_i\Delta$; otherwise, i could profitably deviate by rejecting all proposals in every period. Hence, the set of PBE payoff vectors is a subset of $\{(w_1, \dots, w_n) \in \mathbb{R}^n : \sum_{i=1}^n w_i \leq V^*(\alpha_0) \text{ and } w_i \geq s_i\Delta, \forall i \in N\}$. It follows from Lemma B4 that any PBE that fails to support the optimal stopping rule is Pareto dominated by some renegotiation-proof equilibrium. Therefore, the set of renegotiation-proof equilibrium payoff vectors is $\{(w_1, \dots, w_n) \in \mathbb{R}^n : \sum_{i=1}^n w_i = V^*(\alpha_0) \text{ and } w_i \geq s_i\Delta, \forall i \in N\}$ (Lemma B4).

(ii) Other collegial rules. Suppose first that $\emptyset \neq V \equiv \bigcap \mathcal{D} \notin \mathcal{D}$. Let $m = |N \setminus V|$, and let $\bar{\Delta} \equiv \sup\{\Delta \in \mathbb{R}_+ : 2(1 - e^{-\rho\Delta}) < e^{-\rho\Delta}/m\} > 0$. To establish the result in this case, we will use an analogous argument to that used for noncollegial rules: we will show that every payoff vector in $W \equiv \{(w_1, \dots, w_n) \in \mathbb{R}^n : \sum_{i=1}^n w_i = V^*(\alpha_0), w_i \geq s_i\Delta \forall i \in V, \text{ and } w_i \geq 0 \forall i \notin V\}$ can be supported in a renegotiation-proof equilibrium. As the set of PBE payoff vectors must be contained in $\{(w_1, \dots, w_n) \in \mathbb{R}^n : \sum_{i=1}^n w_i \leq V^*(\alpha_0), w_i \geq s_i\Delta \forall i \in V, \text{ and } w_i \geq 0 \forall i \notin V\}$, this implies that the set of renegotiation-proof PBE vectors is W .

Take an arbitrary $w^0 \in W$, and let $y^0 \in X$ be defined by $y_i^0 \equiv w_i^0/V^*(\alpha_0)$ for all $i \in N$. Our objective is to construct a PBE σ , in which: (i) the committee implements the optimal stopping rule in every continuation game (so that σ is renegotiation-proof); (ii) on

the path, aggregate revenues are distributed according to y^0 in every period. To this end, we first define revenue distributions $y^i(p | \alpha) \in X$, for all $i \in N$ and $p \in \{R, S\} \times [0, 1] \times X$, as follows: Let $S(p | \alpha) \equiv V^*(\alpha) - \sum_{k \in V} w_k(p | \alpha)$; and let

$$y_j^i(p | \alpha) \equiv \begin{cases} \frac{w_i(p | \alpha)}{V^*(\alpha)} & \text{if } j \in V, \\ 0 & \text{if } j = i \text{ \& } j \notin V, \\ \frac{1}{m} S(p | \alpha) & \text{if } j \neq i \in V \text{ \& } j \notin V, \\ \frac{1}{m-1} S(p | \alpha) & \text{if } j \neq i \notin V \text{ \& } j \notin V, \end{cases}$$

for all $j \in N$. As the stopping rule optimal is implemented in every continuation game both on and off the path, such an equilibrium must be renegotiation-proof.

We define the strategy profile σ in terms of “phases,” formally represented by pairs in $\{1, \dots, n\} \times \{0, 1, \dots, n\}$. Every phase (ℓ, i) prescribes behavior in the ℓ th proposal stage of any given period and in the n voting stages that follow it: “ i ” indicates that σ prescribes policy $(a^*(\alpha), 1, y^i)$ to be implemented. Specifically, in any period in where the status quo is p , the belief is $\alpha \in A$ and the order of proposers is $(\pi_1, \dots, \pi_n)$, if the game is in phase $(\ell, i) \neq (n, \pi_n)$, then σ prescribes the following behavior:

(P1) proposer π_ℓ offers policy $(a^*(\alpha), 1, y^i)$;

(V1.a) if π_ℓ offered $(a^*(\alpha), 1, y^i)$ where $i \notin V$, then every voter $j \neq i$ accepts it, and voter i accepts it if and only if one of the following conditions hold: she is the first voter; or all the previous voters accepted $(a^*(\alpha), 1, y^i)$; or some of the previous voters rejected it and

$$y_i^i(p | \alpha) V^*(\alpha) \geq \begin{cases} (1 - \delta) w_i(p | \alpha) + \delta \mathbb{E}[y_i^k(p | \tilde{\alpha}) V^*(\tilde{\alpha}) | \alpha, p] & \text{if } \ell = n, \\ & \text{or } \ell = n - 1 \text{ \& } i = \pi_n, \\ y_i^k(p | \alpha) V^*(\alpha) & \text{otherwise,} \end{cases}$$

where k is the last of the previous voters who rejected it.

(V1.b) if π_ℓ offered $(a^*(\alpha), 1, y^i)$ where $i \in V$, then all voters accept it.

(V2a) if π_ℓ offered any $p' \neq (a^*(\alpha), 1, y^i)$ in $P(p | \alpha) \equiv \{p'' \in \{R, S\} \times [0, 1] \times X : w_j(p'' | \alpha) \leq w_j(p | \alpha), \forall j \in V\}$, then each voter j acts as follows:

- If she is a vetoer and $w_j(p' | \alpha) \leq w_j(p | \alpha)$, then she rejects p' (irrespective of the previous voters' choices);

- otherwise, she accepts p' if and only if $(1 - \delta)w_j(p' | \alpha) + \delta\mathbb{E}[y_j^k(p' | \tilde{\alpha})V^*(\tilde{\alpha}) | \alpha, p]$ is greater than

$$\begin{cases} (1 - \delta)w_j(p | \alpha) + \delta\mathbb{E}[y_j^k(p | \tilde{\alpha})V^*(\tilde{\alpha}) | \alpha, p] & \text{if } \ell = n - 1 \text{ \& } k = \pi_n, \text{ or } \ell = n, \\ y_j^k(p | \alpha)V^*(\alpha) & \text{otherwise,} \end{cases}$$

where k is the last of the vetoers in $V^r \equiv \{i \in V : w_j(p' | \alpha) \leq w_j(p | \alpha)\}$ in the sequence of voters.

(V2b) if π_ℓ offered any $p' \neq (a^*(\alpha), 1, y^i)$ outside $P(p | \alpha)$, then each voter j acts as follows:

- If she is a vetoer, then she accepts p' if and only if $w_j(p' | \alpha) \geq w_j(p | \alpha)$;
- if she is not a vetoer, then she rejects p' .

If the game is in phase (n, π_n) , then σ prescribes the following behavior: (i) proposer π_n passes; and (ii) if π_n proposed some policy $p' \neq p$, then j behaves as in case (V2a) if $p' \in P(p | \alpha)$, and as in case (V2b) otherwise.

Phases evolve according to the following recursive rules. In period 1, play begins in phase $(1, 0)$. Then in every period, at the end of any sequence of votes that began in any phase $(\ell, i) \neq (n, \pi_n)$:¹²

(t1.a) If policy $(a^*(\alpha), 1, y^i)$, where $i \notin V$, was proposed and accepted by all voters but i , then the game transitions to phase $(1, i)$;

(t1.b) If policy $(a^*(\alpha), 1, y^i)$, where $i \in V$, was proposed and accepted by all voters, then the game transitions to phase $(1, i)$;

(t2.a) if policy $(a^*(\alpha), 1, y^i)$, where $i \notin V$, was accepted but some voters different from i rejected it, then the game transitions to phase $(1, k)$, where k is the last of those voters;

(t2.b) if policy $(a^*(\alpha), 1, y^i)$, where $i \in V$, was accepted but some voters rejected it, then the game transitions to phase $(1, k)$, where k is the last of those voters;

(t3.a) if policy $(a^*(\alpha), 1, y^i)$, where $i \notin V$, was proposed and rejected, then the game transitions to phase $(\ell + 1, k)$, where k is the last voter different from i who rejected $(a^*(\alpha), 1, y^i)$;

¹²We set $\ell + 1 = 1$ whenever $\ell = n$.

(t3.b) if policy $(a^*(\alpha), 1, y^i)$, where $i \in V$, was proposed and rejected, then the game transitions to phase $(\ell + 1, k)$, where k is the last voter who rejected $(a^*(\alpha), 1, y^i)$;

(t4) if the status quo differs from $(a^*(\alpha), 1, y^i)$ and proposer π_ℓ passes, then the game moves to phase $(\ell + 1, \pi_\ell)$;

(t5) if proposer π_ℓ offered a policy $p' \neq (a^*(\alpha), 1, y^i)$ in $P(p \mid \alpha)$ and her proposal was rejected by all vetoers j such that $w_j(p' \mid \alpha) \leq w_j(p \mid \alpha)$, then the game transitions to phase $(\ell + 1, \pi_\ell)$;

(t6) if proposer π_ℓ offered a policy $p' \neq (a^*(\alpha), 1, y^i)$ in $P(p \mid \alpha)$ and her proposal was rejected by the committee, but accepted by some vetoers j such that $w_j(p' \mid \alpha) \leq w_j(p \mid \alpha)$, then the game transitions to phase $(\ell + 1, k)$, where k is the last of those vetoers who accepted p' ;

(t7) if proposer π_ℓ offered a policy $p' \neq (a^*(\alpha), 1, y^i)$ in $P(p \mid \alpha)$ and her proposal was accepted by the committee, then the game transitions to phase $(1, k)$, where k is the last vetoer j such that $w_j(p' \mid \alpha) \leq w_j(p \mid \alpha)$;

(t8) if proposer π_ℓ offered a policy $p' \neq (a^*(\alpha), 1, y^i)$ outside $P(p \mid \alpha)$ and her proposal was rejected by all voters in $N \setminus V$, then the game transitions to phase $(\ell + 1, \pi_\ell)$;

(t9) if proposer π_ℓ offered a policy $p' \neq (a^*(\alpha), 1, y^i)$ outside $P(p \mid \alpha)$ and her proposal was rejected by the committee, but accepted by some voters in $N \setminus V$, then the game transitions to phase $(\ell + 1, k)$, where k is the last of the voters in $N \setminus V$ who accepted p' ;

(t10) if proposer π_ℓ offered a policy $p' \neq (a^*(\alpha), 1, y^i)$ outside $P(p \mid \alpha)$ and her proposal was accepted by the committee, then the game transitions to phase $(1, k)$, where k is the last of the voters in $N \setminus V$ who accepted p' ;

At the end of any sequence of votes that began in phase (n, π_n) , we have the following transitions that parallel cases (t5)-(t10) above:

(t11) If the proposer π_n passed, then the game transitions to phase $(1, \pi_n)$;

(t12) if proposer π_ℓ offered a policy $p' \neq p$ in $P(p \mid \alpha)$ and her proposal was rejected by all vetoers j such that $w_j(p' \mid \alpha) \leq w_j(p \mid \alpha)$, then the game transitions to phase $(1, \pi_n)$;

(t13) if proposer π_ℓ offered a policy $p' \neq p$ in $P(p \mid \alpha)$ and her proposal was rejected by the committee, but accepted by some vetoers j such that $w_j(p' \mid \alpha) \leq w_j(p \mid \alpha)$, then

the game transitions to phase $(1, k)$, where k is the last of those vetoers who accepted p' ;

(t14) if proposer π_ℓ offered a policy $p' \neq p$ in $P(p \mid \alpha)$ and her proposal was accepted by the committee, then the game transitions to phase $(1, k)$, where k is the last vetoer j such that $w_j(p' \mid \alpha) \leq w_j(p \mid \alpha)$;

(t15) if proposer π_ℓ offered a policy $p' \neq p$ outside $P(p \mid \alpha)$ and her proposal was rejected by all voters in $N \setminus V$, then the game transitions to phase $(1, \pi_\ell)$;

(t16) if proposer π_ℓ offered a policy $p' \neq p$ outside $P(p \mid \alpha)$ and her proposal was rejected by the committee, but accepted by some voters in $N \setminus V$, then the game transitions to phase $(1, k)$, where k is the last of the voters in $N \setminus V$ who accepted p' ;

(t17) if proposer π_ℓ offered a policy $p' \neq p$ outside $P(p \mid \alpha)$ and her proposal was accepted by the committee, then the game transitions to phase $(1, k)$, where k is the last of the voters in $N \setminus V$ who accepted p' ;

We now verify that for $\Delta < \bar{\Delta}$, this strategy profile is a PBE. We begin with committee member j 's voting behavior. Consider in any period in where the status quo is p , the belief is $\alpha \in A$ and the order of proposers is $(\pi_1, \dots, \pi_n)$. There are several cases:

- *Case 1: In phase $(\ell, i) \neq (n, \pi_n)$, π_ℓ has proposed $(a^*(\alpha), 1, y^i)$.* Observe first that it follows from the definition of σ that voter j 's continuation value at the start of any phase $(\hat{\ell}, \hat{i})$ is

$$\begin{cases} y_j^i(p \mid \alpha) V^*(\alpha) & \text{if } \ell < n, \\ (1 - \delta)w_j(p \mid \alpha) + \delta \mathbb{E}[y_j^i(p \mid \tilde{\alpha}) V^*(\tilde{\alpha}) \mid \alpha, p] & \text{if } \ell = n. \end{cases}$$

We assume that $i \notin V$; the case where $i \in V$ is analogous — just replace “all the previous voters but i ” and “the previous voters different from i ” by “all the previous voters.” We consider several cases in turn:

(1.a) $i = 0$ (so that $\ell = 1$ and $\alpha = \alpha_0$).

(1.a.i) $j \neq i$. Voter j is better off accepting $(R, 1, y^0)$, as prescribed. Indeed, if she is the first voter, or if all the previous voters but i have accepted $(R, 1, y^0)$, then she receives $y_i^0 V^*(\alpha_0)$ if she accepts; while if she rejects, then the game moves to phase $(2, j)$ and she receives $y_j^j(p \mid \alpha_0) V^*(\alpha_0)$; and, by construction $y_i^0 \geq y_j^j(p \mid \alpha_0)$. If some of the previous voters different from i have rejected $(R, 1, y^0)$, then she receives $y_j^j(p \mid \alpha_0) V^*(\alpha_0)$ if she

also rejects it, and $y_j^k(p \mid \alpha_0)V^*(\alpha_0) > y_j^j(p \mid \alpha_0)V^*(\alpha_0)$ if she instead accepts it, where k is the last of the previous voters different from i who rejected it.

(1.a.ii) $j = i$. If voter i is the first voter, or if all the previous voters have accepted $(R, 1, y^0)$, then her choice does not affect her payoff. If some of the previous voters rejected $(R, 1, y^0)$, then her decision can only affect her payoff if she is pivotal. In the latter case, her voting strategy prescribes her to accept if and only if her continuation value from accepting is greater than or equal to her continuation value from rejecting. Hence, she cannot profitably deviate from σ .

(1.b) $i \neq 0$.

(1.b.i) $j \neq i$. There are several cases:

- $\ell < n - 1$, or $\ell = n - 1$ and $j \neq \pi_n$. In this case, the same argument as in (1.a) shows that voter j is better off accepting $(a^*(\alpha), 1, y^i(p \mid \alpha))$, as prescribed by σ .
- $\ell = n - 1$ and $j = \pi_n$. If j is the first voter, or if all the previous voters different from i have accepted $(a^*(\alpha), 1, y^i(p \mid \alpha))$, then her payoff is $y_j^i V^*(\alpha)$ if she also accepts it, and $(1 - \delta)w_j(p \mid \alpha) + \delta \mathbb{E}[y_j^j(p \mid \tilde{\alpha})V^*(\tilde{\alpha}) \mid \alpha, p]$ if she rejects it. We have $y_j^i V^*(\alpha) = (1 - \delta)w_j(p \mid \alpha) + \delta \mathbb{E}[y_j^j(p \mid \tilde{\alpha})V^*(\tilde{\alpha}) \mid \alpha, p] = w_j(p \mid \alpha)$ when j is a vetoer, and

$$\begin{aligned} y_j^i V^*(\alpha) &\geq \frac{1}{m} S(p \mid \alpha) > (1 - \delta) S(p \mid \alpha) \geq (1 - \delta) w_j(p \mid \alpha) \\ &= (1 - \delta) w_j(p \mid \alpha) + \delta \mathbb{E}[y_j^j(p \mid \tilde{\alpha}) V^*(\tilde{\alpha}) \mid \alpha, p] \end{aligned}$$

when she is not a vetoer. (The strict inequality follows from $\Delta < \bar{\Delta}$.) Hence, j cannot profitably deviate from accepting the proposal (as prescribed by σ) in this case.

Now suppose that some of the previous voters different from i have rejected $(a^*(\alpha), 1, y^i(p \mid \alpha))$. If j 's choice is not pivotal, then she is better off accepting the proposal (as prescribed by σ) in order to ensure a transition to a phase where she will receive her largest continuation payoff. If j 's choice is pivotal and she is a vetoer, then she is indifferent between accepting and rejecting: in both cases, some committee member k (possibly equal to j) will be “punished” and she will receive $w_j(p \mid \alpha)$. If j 's choice is pivotal and she is not a vetoer, then she has two options:

- If she votes to accept $p^* \equiv (a^*(\alpha), 1, y^i(p \mid \alpha))$ (so that it is accepted by the commit-

tee), then p^* is implemented and the game transitions to phase $(1, k)$, where $k \neq j$ is the last voter different from i who voted to reject p^* . In this case, j receives a payoff of

$$(1 - \delta)w_j(p^* | \alpha) + \delta \mathbb{E}[y_j^k(p^* | \tilde{\alpha})V^*(\tilde{\alpha}) | \alpha, p^*] = \frac{1}{\tilde{m}}S(p^* | \alpha) = \frac{1}{\tilde{m}}S(p | \alpha) \geq \frac{1}{m}S(p | \alpha) ,$$

where the first equality follows from $y_j^i(p^* | \alpha) = y_j^k(p^* | \alpha)$ (since $j \notin \{i, k\}$), and the second from the fact that $w_\ell(p^* | \alpha) = w_\ell(p | \alpha)$ for all $\ell \in V$.

◦ If she votes to reject $p^* \equiv (a^*(\alpha), 1, y^i(p | \alpha))$ (so that it is rejected by the committee), then the game moves to phase (n, π_n) , in which she will first pass as a proposer and will then receive her “punishment payoff.” That is, she obtains

$$(1 - \delta)w_j(p | \alpha) + \delta \mathbb{E}[y_j^j(p | \tilde{\alpha})V^*(\tilde{\alpha}) | \alpha, p] \leq (1 - \delta)S(p | \alpha) < \frac{1}{m}S(p | \alpha) ,$$

where the first inequality follows from the fact that $y_j^j(p | \tilde{\alpha})$ for all $\tilde{\alpha} \in A$, and the second from $\Delta < \bar{\Delta}$.

We conclude that voter j is better off accepting the proposal, as prescribed from σ .

- $\ell = n$ (so that $i \neq \pi_n$). The argument is exactly the same as in the case where $\ell = n - 1$ and $j = \pi_n$. (In particular, as in that case, if the proposal is rejected both by j and by the committee, then the status quo policy p is implemented and j receives her lowest continuation value from the next period on.)

(1.b.i) $j = i$. The argument is exactly the same as in case (1.a.ii).

- *Case 2: In phase $(\ell, i) \neq (n, \pi_n)$, π_ℓ has proposed a policy $p' \neq (a^*(\alpha), 1, y^i)$ in $P(p | \alpha)$.*
 (2.a) j is a vetoer and $w_j(p' | \alpha) \leq w_j(p | \alpha)$. Suppose first that, given the previous voters' choices and the remaining voters' strategies, voter j 's decision is not pivotal. In this case, her choice only affects the transition to the next phase. It follows from the transition rules (t5)-(t7) that she is always (weakly) better off rejecting p' , as prescribed by σ . Now suppose that her decision is pivotal. If she rejects p' then, from (t5)-(t7), the game transitions to phase $(\ell + 1, k)$ for some $k \neq j$, and she receives a payoff of $(1 - \delta)w_j(p | \alpha) + \delta \mathbb{E}[y_j^k(p | \tilde{\alpha})V^*(\tilde{\alpha}) | \alpha, p] = y_j^k(p | \alpha)V^*(\alpha) = w_j(p | \alpha)$. If instead she deviates, then the game transitions to phase $(\ell + 1, j)$ (see (7)), and she receives $(1 - \delta)w_j(p' | \alpha) + \delta \mathbb{E}[y_j^j(p' | \tilde{\alpha})V^*(\tilde{\alpha}) | \alpha, p] = w_j(p' | \alpha)$. As $w_j(p' | \alpha) \leq w_j(p | \alpha)$, she is better off rejecting.

(2.b) *Either j is not a vetoer or $w_j(p' | \alpha) > w_j(p | \alpha)$.* If, given the previous voters' choices and the remaining voters' strategies, voter j 's decision is not pivotal then, as above, any deviation from σ is unprofitable. If j 's vote is pivotal given the previous voters' moves and the remaining voters' strategies, then it must be the case that all the vetoers in V^r have already voted and they all chose to accept p' . Let k be the last member of V^r who moved. If j chooses to accept p' , then her payoff will be $(1 - \delta)w_j(p' | \alpha) + \delta\mathbb{E}[y_j^k(p' | \tilde{\alpha})V^*(\tilde{\alpha}) | \alpha, p]$; if she chooses to reject p' , then her payoff will be

$$\begin{cases} (1 - \delta)w_j(p | \alpha) + \delta\mathbb{E}[y_j^k(p | \tilde{\alpha})V^*(\tilde{\alpha}) | \alpha, p] & \text{if } \ell = n - 1 \text{ \& } k = \pi_n, \text{ or } \ell = n, \\ y_j^k(p | \alpha)V^*(\alpha) & \text{otherwise.} \end{cases}$$

It follows that j cannot profitably deviate from σ .

- *Case 3: In phase $(\ell, i) \neq (n, \pi_n)$, π_ℓ has proposed a policy $p' \neq (a^*(\alpha), 1, y^i)$ outside $P(p | \alpha)$.*

(3a) *j is a vetoer.* By the same logic as above, she cannot profitably deviate from σ if she is not pivotal (given the previous voters' choices and the remaining voters' strategies). Suppose she is pivotal. If she chooses to accept p' , then she receives $(1 - \delta)w_j(p' | \alpha) + \delta\mathbb{E}[y_j^k(p' | \tilde{\alpha})V^*(\tilde{\alpha}) | \alpha, p] = w_j(p' | \alpha)$, where k is the last of the previous voters who accepted p' — there must be such voter, otherwise j would not be pivotal. If she chooses to reject p' , then she receives $w_j(p | \alpha)$: if $\ell = n - 1$ and $k = \pi_n$, or if $\ell = n$, then she gets $(1 - \delta)w_j(p | \alpha) + \delta\mathbb{E}[y_j^k(p | \tilde{\alpha})V^*(\tilde{\alpha}) | \alpha, p] = w_j(p | \alpha)$; otherwise, she gets $y_j^k(p | \alpha)V^*(\alpha) | \alpha, p] = w_j(p | \alpha)$. It follows that she cannot profitably deviate from σ .

(3b) *j is not a vetoer.* By the same logic as in case (2.a) above, voter j cannot profitably deviate from σ if she is not pivotal (given the previous voters' choices and the remaining voters' strategies). If she is pivotal, then there are several cases:

(3.b.i) $\ell < n - 1$, or $\ell = n - 1$ and $j \neq \pi_n$. If j accepts the proposal, then she receives $(1 - \delta)w_j(p' | \alpha) + \delta\mathbb{E}[y_j^j(p' | \tilde{\alpha})V^*(\tilde{\alpha}) | \alpha, p] = (1 - \delta)w_j(p' | \alpha) \leq (1 - \delta)S(p' | \alpha) < (1 - \delta)S(p | \alpha)$, where the equality follows from the fact that $y_j^j(p' | \tilde{\alpha}) = 0$ for all $\tilde{\alpha} \in A$, and the second inequality follows from $p' \notin P(p | \alpha)$ (and, therefore, $\sum_{l \in V} w_l(p' | \alpha) > \sum_{l \in V} w_l(p | \alpha)$). If she rejects p' (as prescribed by σ), then she receives $y_j^k(p | \alpha)V^*(\alpha) = \frac{1}{m}S(p | \alpha) \geq \frac{1}{m}S(p | \alpha) > (1 - \delta)S(p | \alpha)$, where the last inequality follows from $\Delta < \overline{\Delta}$.

(3.b.ii) $\ell = n - 1$ and $j = \pi_n$, or $\ell = n$. If j accepts the proposal then, by the same logic as above, she receives a payoff that is smaller than $(1 - \delta)S(p \mid \alpha)$. If she rejects p' (as prescribed by σ), then she receives $(1 - \delta)w_j(p \mid \alpha) + \delta \mathbb{E}[y_j^k(p \mid \tilde{\alpha})V^*(\tilde{\alpha}) \mid \alpha, p] = (1 - \delta)w_j(p \mid \alpha) + \delta [\frac{1}{\tilde{m}}S(p \mid \alpha) - (1 - \delta)w_j(a^*(\alpha), 1, y_j^k(p \mid \alpha) \mid \alpha)] \geq \delta [\frac{1}{\tilde{m}}S(p \mid \alpha) - (1 - \delta)S(p \mid \alpha)] > (1 - \delta)S(p \mid \alpha)$, where the last inequality follows from $\Delta < \bar{\Delta}$.

- *Case 4: In phase (n, π_n) , π_n has proposed a policy $p' \neq p$.* One can show that voter j cannot profitably deviate from σ by using the same arguments as in Cases 2 and 3.

We now turn to committee member i 's proposal behavior. Consider in any period in where the status quo is p , the belief is $\alpha \in A$ and the order of proposers is $(\pi_1, \dots, \pi_n)$. In phase $(\ell, i) \neq (n, \pi_n)$, σ prescribes her to propose $(a^*(\alpha), 1, y^i)$, thus receiving a payoff of $y_j^i(p \mid \alpha)V^*(\alpha)$. Suppose that either $\ell < n - 1$, or $\ell = n - 1$ and $i \neq \pi_n$. If j deviates by making any proposal $p' \neq (a^*(\alpha), 1, y^i)$, then her proposal is rejected by the committee, the game transitions to phase $(\ell + 1, j)$, and she receives $y_j^j(p \mid \alpha)V^*(\alpha) \leq y_j^i(p \mid \alpha)V^*(\alpha)$. The deviation is therefore unprofitable. Now suppose that either $\ell = n - 1$ and $i = \pi_n$, or $\ell = n$ (so that $j \neq i$). In this case, if j deviates by making any proposal $p' \neq (a^*(\alpha), 1, y^i)$, then her proposal is rejected, the status quo policy p is implemented, and the game transitions to the next period in phase $(1, j)$. Hence, her payoff is equal to $w^d \equiv (1 - \delta)w_j(p \mid \alpha) + \delta \mathbb{E}[y_j^j(p \mid \tilde{\alpha})V^*(\tilde{\alpha}) \mid \alpha, p]$. If j is a vetoer, then $w^d = w_j(p \mid \alpha) \leq y_j^i(p \mid \alpha)V^*(\alpha)$ (which holds with equality whenever $i \neq 0$), and the deviation is therefore unprofitable. If j is not a vetoer, then

$$w^d = (1 - \delta)w_j(p \mid \alpha) \leq (1 - \delta)S(p \mid \alpha) \leq \frac{1}{\tilde{m}}S(p \mid \alpha) = y_j^i(p \mid \alpha)V^*(\alpha) ,$$

where the second inequality follows from $\Delta < \bar{\Delta}$, and the last equality from $i \neq j \notin V$.

Finally, in phase (n, π_n) , σ prescribes proposer $j = \pi_n$ to pass. If she does so, then the status quo policy p is implemented and the period starts in phase $(1, \pi_n)$ (see (t11)). If she deviates, proposing any policy $p' \neq p$, then her proposal is rejected by the committee: if $p' \in P(p \mid \alpha)$, then it is rejected by at least one vetoer (see (V2a)); if $p' \notin P(p \mid \alpha)$, then it is rejected by all the voters without a veto (see V2b). Therefore, the status quo policy is implemented and, from (t12) and (t15), the next period starts in phase $(1, \pi_n)$. It follows

that the proposer receive the same payoff irrespective of her choice and, consequently, cannot profitably deviate from σ .

We now turn to the case where $\emptyset \neq V \equiv \bigcap \mathcal{D} \in \mathcal{D}$, i.e., the voting rule is oligarchich. If $|V| = 1$, then the result is trivial: in every continuation game, the dictator redistributes all revenues to herself and, therefore, has the same objective function as the social planner. If $|V| \geq 3$, then one can use the same equilibrium constructions as above — or as in the $V = N$ case (section B in the main text) — to establish that $\{(w_1, \dots, w_n) \in \mathbb{R}_+^n : \sum_{i=1}^n w_i = V^*(\alpha_0) \text{ and } w_i \geq s_i \Delta, \forall i \in V\}$ is a subset of renegotiation-proof equilibrium payoff vectors. As the set of PBE payoff vectors must be a subset of $\{(w_1, \dots, w_n) \in \mathbb{R}_+^n : \sum_{i=1}^n w_i \leq V^*(\alpha_0) \text{ and } w_i \geq s_i \Delta, \forall i \in V\}$, this proves that the set of renegotiation-proof equilibrium vectors is $\{(w_1, \dots, w_n) \in \mathbb{R}_+^n : \sum_{i=1}^n w_i = V^*(\alpha_0) \text{ and } w_i \geq s_i \Delta, \forall i \in V\}$.

To complete the proof of the result for oligarchic rules, it remains to show that the same is true if $|V| = 2$. In the previous cases, we could sustain PBEs in which a vetoer i would receive a payoff of $w_i(p \mid \alpha)$ in any continuation game $\Gamma(p \mid \alpha)$ because it was always possible to ensure that every proposal giving her more than $w_i(p \mid \alpha)$ would be rejected by at least one decisive coalition. In those equilibria, it was always impossible for i to offer *all* of the decisive voters more than their rewards for rejecting i 's proposals. The only reason why the previous constructions do not apply in the $|V| = 2$ case is that i only needs *one* of the other players (the other vetoer) to accept her proposal. As $\delta < 1$, whenever i is the last proposer (an event that occurs with positive probability in every period), she can always make proposal that gives her more than $w_i(p \mid \alpha)$ and the other vetoer more than the maximum the latter would get by rejecting the proposal. This in turn implies that, in any period where she proposes last, vetoer i can guarantee herself some payoff greater than $w_i(p \mid \alpha)$ by rejecting the first $(n - 1)$ proposals. One must therefore change the lower bounds on the vetoers' continuation values to account for this possibility. As in the previous construction, if a vetoer i deviates then, from the next period on, we “punish” her by giving the other vetoer the total surplus minus i 's minimum continuation value. For each $i \in V$, let ρ_i be the probability is the last proposer in each period, and let $V_i^+(p \mid \alpha)$

and $V_i^-(p \mid \alpha)$ be the solutions to the following functional equations:

$$\begin{aligned} V_i^-(p \mid \alpha) &= (1 - \delta)w_i(p \mid \alpha) + \delta \mathbb{E}[\varpi_i(p \mid \tilde{\alpha}) \mid \alpha, p] \\ V_i^+(p \mid \alpha) &= V^*(\alpha) - (1 - \delta)w_j(p \mid \alpha) - \delta \mathbb{E}[V^*(\tilde{\alpha}) - \varpi_i(p \mid \tilde{\alpha}) \mid \alpha, p] , \end{aligned}$$

where $\varpi_i(p \mid \alpha) \equiv \rho_i V_i^+(p \mid \alpha) + (1 - \rho_i) V_i^-(p \mid \alpha)$ and $j \in V \setminus \{i\}$, for all $p \in \{R, S\} \times [0, 1] \times X$. Intuitively, $V_i^+(p \mid \alpha)$ [resp. $V_i^-(p \mid \alpha)$] stands for vetoer i 's minimum payoff in continuation game $\Gamma(p \mid \alpha)$ conditional on her being the last proposer [resp. not being the last proposer], and $\varpi_i(p \mid \alpha)$ stands for her minimum payoff in continuation game $\Gamma(p \mid \alpha)$ (computed before the realization of the order of proposers). We then obtain the result with an equilibrium construction as above, but substituting $\varpi_i(p \mid \alpha)$ to $w_i(p \mid \alpha)$ for each $i \in V$.

C Proof of Lemma C1

Let $i, j \in N$ with $i \in C^j$, and all $k \in \mathbb{N}$. By definition of W_i , we have

$$\begin{aligned} \frac{W_i^j(1) - (1 - \delta)\gamma\Delta\bar{r} - \delta W_i^0(1)}{\Delta} &= (1 - \delta)\gamma(1 - \hat{\tau})(r_i - \bar{r}) \\ &\quad + \gamma\hat{\tau}\bar{r}[\delta(x_i^j - \bar{x}_i) - (1 - \delta)(1 - x_i^j)] . \end{aligned}$$

As $x_i^j - \bar{x}_i > 0$, there exists $\hat{\Delta}_{i,j}^1 > 0$ such that $W_i^j(1) - (1 - \delta)\gamma\bar{r} - \delta W_i^0(1) > 0$ whenever $\Delta < \hat{\Delta}_{i,j}^1$. By the same logic, if $k \geq k^*$, then there exists $\hat{\Delta}_{i,j}^2 > 0$ such that $W_i^j(\alpha_k) - (1 - \delta)\bar{s}\Delta - \delta W_i^0(\alpha_k) = \bar{s}\Delta[\delta(x_i^j - \bar{x}_i) - (1 - \delta)(1 - x_i^j)] > 0$ whenever $\Delta < \hat{\Delta}_{i,j}^2$. Now suppose that $k < k^*$. Observe that

$$\begin{aligned} \frac{W_i^j(\alpha_k) - W_i^0(\alpha_k)}{\Delta} &= \left\{ [1 - \delta^{k^* - k}(1 - \gamma\Delta)^{k^* - k}] \alpha_k \gamma \bar{r} + \delta^{k^* - k} [1 - \alpha_k + (1 - \gamma\Delta)^{k^* - k} \alpha_k] \bar{s} \right\} \\ &\quad \times (x_i^j - \bar{x}_i) , \end{aligned}$$

where the first bracketed term on the right-hand side represents the expected social welfare (divided by Δ) under the optimal stopping rule. As $\alpha_k > \alpha^*$, this term is greater than or equal to $\bar{s}$. Hence, $W_i^j(\alpha_k) - W_i^0(\alpha_k) \geq \bar{s}\Delta(x_i^j - \bar{x}_i) > 0$, and

$$\frac{W_i^j(\alpha_k) - (1 - \delta)\bar{s}\Delta - \delta W_i^0(\alpha_k)}{\Delta} \geq \frac{1 - \delta}{\Delta} W_i^j(\alpha_k) - (1 - \delta)\bar{s} + \frac{\delta}{\Delta} \bar{s} \hat{\tau} (x_i^j - \bar{x}_i) .$$

An application of l'Hôpital's rule shows that $(1 - \delta)/\Delta \rightarrow \rho$ as $\Delta \rightarrow 0$. As $W_i^j(\cdot)$ and $W_i^0(\cdot)$ are bounded, there exists $\hat{\Delta}_{i,j}^3 > 0$ such that $W_i^j(\alpha_k) > (1 - \delta)\bar{s}\Delta - \delta W_i^0(\alpha_k) > 0$ whenever $\Delta < \hat{\Delta}_{i,j}^3$.

Consider now the last inequality in the lemma. Let $\Psi(\alpha_k) \equiv W_i^j(\alpha_k) - (1 - \delta)\alpha_k\gamma\Delta\bar{r} - \delta\alpha_k\gamma\Delta W_i^0(1) - \delta(1 - \alpha_k\gamma\Delta)W_i^0(\alpha_{k+1})$. Suppose first that $k \geq k^*$. It is readily checked that

$$\lim_{\Delta \rightarrow 0} \frac{\Psi(\alpha_k)}{\Delta} = (1 - \hat{\tau})s_i + \hat{\tau}x_i^j\bar{s} - [(1 - \hat{\tau})s_i + \hat{\tau}\bar{x}_i\bar{s}] = \hat{\tau}(x_i^j - \bar{x}_i)\bar{s} > 0 .$$

Therefore, there exists $\hat{\Delta}_{i,j}^4 > 0$ such that $W_i^j(\alpha_k) > (1 - \delta)\alpha_k\gamma\bar{r} + \delta\alpha_k\gamma\Delta W_i^0(1) + \delta(1 - \alpha_k\gamma\Delta)W_i^0(\alpha_{k+1})$ whenever $\Delta < \hat{\Delta}_{i,j}^4$. Finally, suppose that $k < k^*$. By definition of W_i^j , we have

$$\begin{aligned} \Psi(\alpha_k) &= -(1 - \delta)\alpha_k\gamma\Delta\bar{r}(1 - x_i^j) + \delta\alpha_k\gamma\Delta[W_i^j(1) - W_i^0(1)] \\ &\quad + \delta(1 - \alpha_k\gamma\Delta)[W_i^j(\alpha_{k+1}) - W_i^0(\alpha_{k+1})] \\ &\geq -(1 - \delta)\alpha_k\gamma\Delta\hat{\tau}(1 - x_i^j)\bar{r} + \delta(x_i^j - \bar{x}_i)[\alpha_k\gamma\Delta\hat{\tau}\bar{r} + (1 - \alpha_k\gamma\Delta)\hat{\tau}\bar{s}] \\ &> -(1 - \delta)\alpha_k\gamma\Delta\hat{\tau}(1 - x_i^j)\bar{r} + \delta\hat{\tau}(x_i^j - \bar{x}_i)\bar{s} , \end{aligned}$$

where the first inequality follows from $W_i^j(\alpha_{k+1}) - W_i^0(\alpha_{k+1}) \geq \bar{s}\hat{\tau}(x_i^j - \bar{x}_i)$ (as established above), and the second follows from our assumption that $\gamma\bar{r} > \bar{s}$. Therefore, there exists $\hat{\Delta}_{i,j}^5 > 0$ such that $W_i^j(\alpha_k) > (1 - \delta)\alpha_k\gamma\Delta\bar{r} + \delta\alpha_k\gamma\Delta W_i^0(1) + \delta(1 - \alpha_k\gamma\Delta)W_i^0(\alpha_{k+1})$ whenever $\Delta < \hat{\Delta}_{i,j}^5$. Setting $\bar{\Delta} \equiv \min\{\hat{\Delta}_{i,j}^\ell : i, j \in N \text{ \& } \ell = 1, \dots, 5\}$, we obtain the lemma.