

Two-Party Competition with Persistent Policies*

Jean Guillaume Forand[†]

October 11, 2009

Abstract

This paper studies the Markov perfect equilibrium outcomes of a dynamic game of electoral competition between two policy-motivated parties. I model *incumbent policy persistence*: parties commit to implement a policy for their full tenure in office, and hence in any election only the opposition party renews its platform. In equilibrium, parties alternate in power and policies converge to symmetric alternations about the median voter's ideal policy. Parties' disutility from opponents' policies leads to alternations that display *bounded extremism*; alternations far from the median are never limits of equilibrium dynamics. Under a natural restriction on strategies, I find that *robust* long-run outcomes display *bounded moderation*; alternations close to the median are reached in equilibrium only if policy dynamics start there.

JEL Classification: C73, D72, D78

Keywords: Dynamic Electoral Competition, Policy Persistence, Policy Convergence, Markov Perfect Equilibrium

1 Introduction

If the parties are viewed in [a] temporal framework, one may better appraise the old saw that the parties offer the electorate only a choice between tweedledum and tweedledee. In fact, the differences between the parties vary from stage to stage in the conversion of controversy into new consensus.
(Key (1958), p.247.)

Political parties are long-lived organisations that view single elections not as isolated contests but as events in histories of partisan competition. Adopting a dynamic framework can lead to new insights into the policies that parties champion when vying for power and implement once in office. While a wealth of static models help us understand the forces which lead to policy convergence or divergence in one-shot elections,¹ they cannot, by their nature, properly explain

*I would like to thank Li Hao, Martin Osborne and Colin Stewart for their supervision, comments and suggestions. I would also like to thank participants at the NASM 2009 and the CEA 2009 conferences.

[†]Department of Economics, University of Toronto, 150 St. George Street, Toronto, Ontario, M5S 3G7. jg.forand@utoronto.ca

¹See Osborne (1995) for a survey.

the feature V.O. Key refers to above: party competition tends to lead over time to gradual transitions from divergent to convergent outcomes. In this paper, I present a dynamic model to address why in party systems with two dominant parties (i) governments alternate, (ii) starting from differentiated ideological positions, parties gradually move towards proposing platforms which resemble one another yet (iii) they never become as indistinguishable as tweedledum and tweedledee; party labels matter and parties maintain distinct policy goals.

I formulate a dynamic game of electoral competition between two ideological parties that have ideal (single-dimensional) policies on each side of that of the median voter. Voters are myopic and support the party whose current policy yields them higher utility. Incumbency status, which is inherited from previous elections, generates the key dynamic linkage of the model. Under *incumbent policy persistence*, parties commit to enact specific policies if elected for their entire *tenure* in office, as opposed to their current *term*. In each election, incumbents face the voters championing (or rather defending) the policies they implemented in their previous term, while opposition parties are free to choose a new platform. Opposition parties are forward-looking and understand that the platforms that carry them to office will support their bids for reelection. The key insights of my model make precise how opposition parties trade off winning current elections with policies they prefer against committing to more moderate policies in order to constrain their future opponents. Parties are restricted to Markov strategies, which depend on the outcomes of previous elections only insofar as these affect the state: the identity of the incumbent party and its policy.

While the model admits a complex set of Markov perfect equilibria, I characterise its *long-run policy outcomes*, the limit points of equilibrium paths given some initial state. I show that all equilibria have (i) alternation in power and (ii) bounded extremism in the long-run, while robust equilibria have (iii) bounded moderation. From initial states that are sufficiently distant from the median, two-party competition always leads to some convergence. The bound on long-run extremism, which is driven by parties' incentives to impose moderation on their future opponents, is tight. In particular, the indefinite repetition of the median policy can occur in the long-run. However, median convergence is not a *robust* outcome of the model. Under a natural restriction on equilibrium strategies, I show that alternations close to the median occur in the long-run only if policy dynamics start there. That is, while convergence *towards* the median is dynamically robust, convergence *to* the median is not and ideological differentiation is persistent. The incentives to sustain convergence unravel as policies approach the median. The benefit of committing to more moderate policies is that future opponents commit to even more moderate policies, while its cost consists of foregone policy gains in the current election. As parties champion ever more similar policies, discounting wipes out the benefits of policy convergence. Lastly, the bound on robust long-run moderation is tight.

Generating novel policy dynamics and insights that go beyond those of static models requires tangible dynamic linkages between successive elections. Proposition 1 shows that dropping in-

cumbent policy persistence from my model yields a repeated game with a unique subgame perfect equilibrium in which both parties commit to the median policy after all histories. More dramatically, Duggan and Fey (2006) show that *any* policy path can be enforced by some subgame perfect equilibrium of the repeated two-party Downsian model with forward-looking voters. Instead of considering sequences of strategically identical elections, I exploit opposition parties' greater freedom to propose significant shifts in policy. Incumbents are associated to current policies for a variety of reasons: renouncing previous commitments or admitting policy mistakes have large electoral costs; reelected politicians' ideologies, which drive their policy choices, rarely change substantially between terms; voters disregard incumbents' promises of policy change through retrospective voting. In brief, while challengers are evaluated on their *promises*, incumbents are evaluated on their *records*.² The assumption of full commitment is stark but it allows a simple characterisation of equilibrium outcomes. In Section 5.2, I discuss my paper's relationship with those on dynamic legislative bargaining and show how my results persist in a model in which incumbents can revise their policies with positive probability.³ Two features of my setup are critical for my results. First, incumbents understand that they will be evaluated on their records and that choosing non-median policies puts them at a disadvantage relative to their opponents. Second, this disadvantage is larger for incumbent policies that are further away from the median.

My model supports two distinct observed patterns of power and alternation. Proposition 2 shows that on any equilibrium policy paths either the initial policy is absorbing and the incumbent remains in power forever, or power changes hands each election and the dynamics converge to an alternation at policies symmetric about the median. These pairs of policies are the *alternating long-run policy outcomes* of the model. The first type of policy path occurs only if a leftist (rightist) incumbent party is implementing a policy to the right (left) of the median in the initial state, sapping the competitive incentives of its opposition. Otherwise, the party system is competitive, both parties hold office and successive opposition parties win elections by committing to increasingly moderate policies. In the long run, incumbents are defeated by opposition parties that are equally preferred by the median voter.

Few previous dynamic models of elections generate plausible patterns of alternations that persist in the long-run. Predictable left-right alternation is not an essential feature of my results. In fact, the extension in Section 5.2 noted above generates random alternation and staggered convergence paths. Models of dynamic elections with imperfect information about politicians' preferences based on Duggan (2000) typically do not generate alternation in the long-run; successive incumbents survive for one term in office until a sufficiently moderate candidate is elected and survives all challenges. A notable exception is Kalandrakis (2008), who allows incumbents'

²This has been documented by Miller and Wattenberg (1985) and Nadeau and Lewis-Beck (2001) for presidential elections in the United States. Both find evidence that voters tend to evaluate incumbents retrospectively and challengers along prospective lines.

³For dynamic legislative bargaining models, see Baron (1996), Baron, Diermeier and Fong (2008), Bowen and Zahran (2009), Duggan and Kalandrakis (2009), Fong (2008) and Kalandrakis (2004, 2007).

preference types to be redrawn. Hence an incumbent previously believed to be moderate may suddenly choose to implement extreme policies and be replaced. In Van Weelden (2009), while preferences are commonly known, there is no alternation on the equilibrium path as voters use the threat of alternation to induce candidates to enact moderate policies.

Proposition 2 also shows that, even with strong initial ideological disagreement, parties' policies gradually become more moderate over time and tend to reach a stable yet differentiated state. In the United States, periods separating what Key (1955) has termed 'critical elections' have been shown to consist of a process of stabilisation in which 'polarization gives way to conciliation. As it does, the parties move from the poles toward the center and the distance between them narrows.'⁴ Another well-known formulation of electoral competition's tendency towards stabilisation is Lipset and Rokkan's (1967) 'freezing' hypothesis, which claims that 'the party systems of the 1960's reflect [...] the cleavage structures of the 1920's.'⁵

Proposition 3 characterises the set of alternating long-run policy outcomes, which consists of all sufficiently moderate policy alternations. I show that a tight upper bound on the extremism of any alternating outcome reached in the long-run is given by the *most moderate* of the parties' preferred alternations. Discounting ensures that both parties prefer some alternation to the repetition of the median, as when alternations are sufficiently moderate the gain from enacting policies on their side of the median dominates the discounted disutility of opponents' policies. This bound on long-run extremism makes precise the implications of the model's key dynamic trade-off by identifying the policies that are sufficiently extreme to provide incentives for some party to unilaterally enact a more moderate policy in order to rein in its future opponents. That the bound on long-run extremism is tight follows from equilibrium construction.

Of the long-run outcomes of the model, some are reached only if they occur in the initial state, whereas others can be reached from more extreme states through sequences of elections decided by increasingly moderate policies. A *robust long-run policy outcome* is a long-run policy outcome that can be reached from some initial state with a policy that differs from the policy outcome itself. To study robust outcomes, I require that parties' strategies be *consistent*. Such Markov strategies do not allow parties to condition on an incumbent's exact policy when committing to policies in the interior of their set of winning policies. Since incumbents' policies matter to opposition parties only when they constrain their policy choices, consistency strengthens the requirement that Markov strategies depend only on payoff-relevant information. Consistency is a natural restriction and it allows simple characterisations of payoffs and policies on equilibrium

⁴Sundquist (1983), p.319. Key (1955) defines critical elections as 'a type of election in which there occurs a sharp and durable electoral realignment between parties' (p.16). A oft-cited example is the presidential election of 1932 that brought F.D. Roosevelt and the New Deal to power. Documenting the aftereffects of that election, Sundquist (1983) notes that 'as the polarization of the electorate that had characterised the depression years dissolved into the moderation of more prosperous times, the conflict between the parties was somewhat muted. They remained anchored on either side of the activist-conservative line of cleavage, but the distance between them that had been so great in the early 1930's diminished.' (p.337)

⁵Lipset and Rokkan (1967), p.50.

convergence paths, which are otherwise complex.

Proposition 4 characterises the set of robust long-run policy outcomes under equilibria in consistent strategies and shows it to contain all alternating outcomes that are sufficiently extreme. This tight bound on the moderation of robust long-run outcomes is derived explicitly and is strictly away from the median. Alternating outcomes close to the median are never reached by consistent equilibrium policy dynamics that start from more extreme states; they are long-run outcomes only if policy dynamics start there. On an equilibrium convergence path, moderate policy commitments are supported by opponents' promises of further moderate commitments in future elections. That is, moderation must be self-reinforcing. The proof that robust long-run outcomes have bounded moderation shows that the incentives to commit to moderate policies unravel as convergence paths approach the median. In particular, I construct bounds on how much policy moderation each party is willing to implement at each step of a convergence path in response to an opponent's proposed moderate move in the next election. When policy dynamics are sufficiently close to the median, parties' 'demands' for moderation are incompatible and neither party has sufficient incentive to commit to more moderate policies and sustain convergence. Discounting is critical to this argument. As policies approach the median, moderate moves of similar size by a party and its opponent have similar effects (in absolute value) on its payoffs, yet a party suffers the loss from its moderate move in full today while the gain from its opponent's moderation is discounted. The bound on robust long-run moderation is shown to be tight through equilibrium construction.

Dynamic models of asynchronous policy competition can be traced back to Downs (1957) and were first formally presented in Kramer (1977) and Wittman (1977). They study models similar to mine in which, crucially, parties are myopic. Their models differ from each other only in their assumptions about parties' preferences. Kramer (1977) assumes that parties are office-motivated and maximise votes, while Wittman (1977) assumes that parties are policy-motivated.⁶

More generally, the idea that forward-looking incumbents have incentives to strategically position current policies to affect future political outcomes has had numerous applications.⁷ More closely linked to my paper are the infinite horizon models of dynamic legislative bargaining, spatial electoral competition and public goods provision. In dynamic legislative bargaining models, a legislator is recognised each period to propose some policy which is put to a vote against the status quo. If it passes, it is implemented and becomes next period's status quo. As opposed to my characterisation of equilibrium outcomes, papers on dynamic legislative bargaining typically

⁶Related to these papers is the literature on competition between myopic adaptive parties, such as Kollman, Miller and Page (1992,1998), de Marchi (1999) and Laver (2005). Kollman, Miller and Page (1992) generate policy dynamics that moderate over time yet stay bounded away from the median in the long-run. In their framework, these dynamics result from policy experimentation by myopic parties.

⁷In a well-known paper, Alesina and Tabellini (1990) show how incumbents accumulate excessive public debt in order to 'tie the hands' of future governments that may not share their preferences over public goods spending. For a review of this literature, consult Persson and Tabellini (2000). Bai and Lagunoff (2009) also present a useful discussion of this literature in the context of their more general model.

study specific equilibria. Baron’s (1996) model is most closely related to mine. He characterises an equilibrium in which all policy paths converge to the median policy, which contrasts with the non-robustness of policy outcomes near the median in my model. His result follows from the median legislator eventually being recognised, proposing the median policy and never supporting anything other than the status quo in future periods. In Section 5.2, I show that it is Baron’s (1996) assumption of the existence of a median legislator that is critical for median convergence. In the legislative bargaining version of my model, robust convergence outcomes are still bounded away from the median.

Dynamic models of electoral competition between candidates that are privately informed about their policy preferences generate incentives to choose moderate policies to maintain a reputation for moderate preferences in order to reap the benefits power as well as keep opponents from implementing their own policies.⁸ In these models, candidate selection by parties is non-strategic and candidates’ informational advantage is derived from having been drawn at random from the voting population or the party’s membership,⁹ while in my model parties can commit to any policy. In the absence of signalling by privately informed candidates, Van Weelden (2009) shows that similar intuition and dynamics can obtain. However, as noted above, these models typically do not generate alternation in the long-run as all sufficiently moderate candidates are maintained in power. In contrast, all incumbents are replaced on most equilibrium paths of my model; moderation does not guarantee reelection, since opponents can respond by championing more moderate policies themselves.

Less closely related, Azzimonti (2009) embeds two-party probabilistic elections into a dynamic growth model in which two segments of the population differ in their taste for the composition of public goods spending, and studies the inefficiencies in agents’ capital accumulation decisions due to political instability.¹⁰ Bai and Lagunoff (2009) present a general dynamic model of ‘policy-endogenous political power’ in which current policies determine future political power and provide an application to the provision of public goods. Acemoglu, Egorov and Sonin (2009) study the long-run outcomes of a dynamic constitutional choice game which they call dynamically stable states: the states that persist in the long-run. In contrast to my results, their equilibrium characterisations rely on players being arbitrarily patient.

⁸See Banks and Duggan (2006), Bernhardt, Duggan and Squintani (2007), Bernhardt, Campuzano, Squintani and Câmara (2009), Duggan (2000) and Kalandrakis (2008).

⁹Bernhardt, Campuzano, Squintani and Câmara (2009), show that drawing opponents from opposite sides of the political spectrum (i.e., from different parties) makes incumbents more willing to compromise by lowering their continuation value if they lose office. Related my assumption of incumbent policy persistence, Kalandrakis (2008) assumes that a party that a party that has recently lost an election is more likely to field a candidate of a different preference type.

¹⁰Battaglini and Coate (1997,1998) present related models in which political outcomes are determined by legislative bargaining.

2 Model

Two parties, L and R , contest an infinite sequence of elections at times $t = 0, 1, \dots$. Each period starts with the incumbent party $I \in \{L, R\}$ in power, and the remaining party in opposition. An election consists of a vote over which party should form the next government, with the winning party determined by majority rule. The opposition party $-I = \{L, R\} \setminus \{I\}$ commits to implementing a policy in the policy space $X = [0, 1]$, if elected, and for as long as it remains in power: this is the assumption of incumbent policy persistence. Hence, in any election, the incumbent's policy commitment is inherited from the election that brought it to power.¹¹ A party may also choose not to participate in the election.

A large (odd) number of voters have single-peaked preferences over policies. Their ideal policies are distributed over policy space X and some voter's ideal policy M corresponds to the median of voters' ideal policies. It is without loss of generality to assume that the median voter has distance preferences. Voters are myopic and in all voting subgames, I restrict attention to the equilibrium in weakly undominated strategies in which voters support the party that will enact a policy closest to their ideal policy if brought to power in this election. Hence the median voter is decisive and the party whose policy is closest to M wins the election. I assume for simplicity that ties are broken in favour of the opposition party. Myopic voting has the benefit of focusing attention solely on the competition between the parties. However, in Section 5.1, I show that all the equilibrium outcomes studied in the paper persist if voters are forward-looking.

To formalise the dynamic game, define a *state* (I, x) , with $I \in \{L, R\}$ and $x \in X$, which records the identity of the incumbent party along with its policy commitment. Given a state (I, x) , the corresponding stage game is a single-agent decision problem with the following timing:

- The opposition party $-I$ commits to a policy $z \in X$, or does not contest the election, written $z = \{Out\}$.
- Elections are held. Party I wins if and only if $|x - M| < |z - M|$.
- Parties have single-peaked preferences over policies represented by u_L and u_R , with peaks at policies 0 and 1 respectively. Suppose, without loss of generality, that $M \leq \frac{1}{2}$, so that party L is (weakly) favoured by the median voter. It is not critical that parties' ideal policies are located at the extremes of the policy space, only that they be on each side of M . Assume that $u_L(0) = u_R(1) = 0$, $u_L(u_R)$ is strictly decreasing (increasing), twice

¹¹A natural interpretation of the assumption of incumbent policy persistence is to view parties as consisting of party leaders and a large number of members with the leaders of the opposition party selecting a candidate from their membership to represent the party in the upcoming elections (on this, see Bernhardt, Campuzano, Squintani and Câmara (2009) and Kalandrakis (2009)). As in citizen-candidate models (Osborne and Slivinski (1996), Besley and Coate (1997)), candidates have known preferences over policies and enact their ideal policy if leading a government. While parties that lose elections can replace their representatives, the leaders of the governing party cannot change their party's candidate, and that party contests the current election with the candidate that won the previous contest.

continuously differentiable and strictly concave.¹²

Given state (I, x) , let $W(I, x)$ be the *set of winning policies* for the opposition party. Note that for any $x \in X$, $W(R, x) = W(L, x) = [\min\{2M - x, x\}, \max\{2M - x, x\}]$ and $W(R, x') \subset W(R, x)$ and $W(L, x') \subset W(L, x)$ whenever $|x' - M| \leq |x - M|$. Payoffs to parties and the median voter, along with the set of winning policies $W(L, x)$ for some policy x , are illustrated in Figure 1.

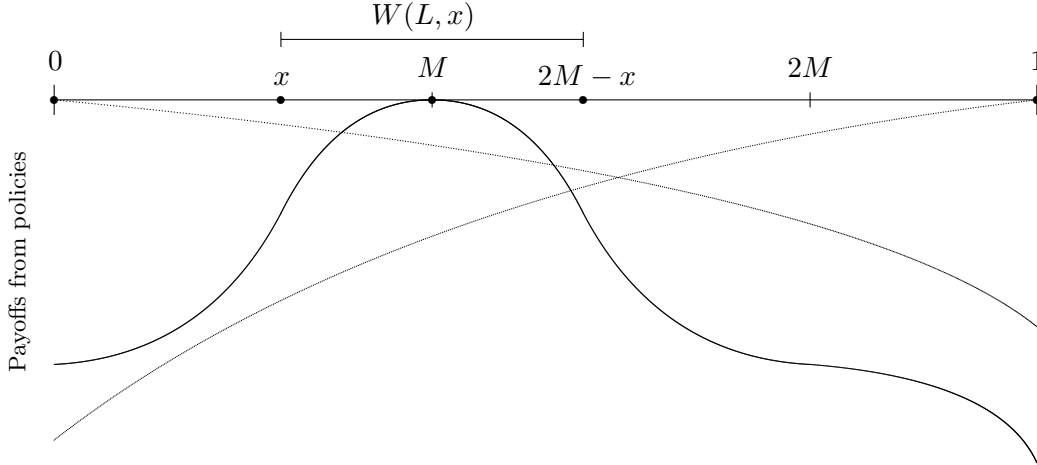


Figure 1: Policy Space and Parties' Preferences.

Transitions between states are given as follows: the current period's winning party and policy become next period's incumbent party and incumbent policy, respectively. Formally, define the *state transition function* $\tau : (\{L, R\} \times X) \times (X \cup \{Out\}) \rightarrow \{L, R\} \times X$ by

$$\tau((I, x), z) = \begin{cases} (I, x) & \text{if } |x - M| < |z - M| \text{ or } z = Out, \\ (-I, z) & \text{if } |x - M| \geq |z - M|. \end{cases}$$

The dynamic game proceeds as follows. Given some initial state (I^0, x^0) , the two parties take part in an infinite sequence of elections, where the transition between stage games is given by τ . A *history starting from* (I^0, x^0) is a sequence $\{(I^i, x^i)\}_{i=0}^N \in (\{L, R\} \times X)^N$ with $N \leq \infty$ such that for $i \geq 1$, $(I^i, x^i) = \tau((I^{i-1}, x^{i-1}), z)$ for some $z \in X \cup \{Out\}$. The payoff to party J from terminal history $\{(I^i, x^i)\}_{i=0}^\infty$ starting from (I^0, x^0) is

$$\sum_{i=1}^{\infty} \delta_J^{i-1} u_J(x^i),$$

¹²Concavity simplifies the results but can be relaxed. See Section 4.2.

where $\delta_J < 1$ is party J 's discount factor.

Definition 1. A *Markov strategy* for party J is a function $\sigma_J : \{L, R\} \times X \rightarrow X \cup \{Out\}$, with the restriction that $\sigma_J(J, x) = x$ for all $x \in X$.

The restriction captures the assumption of incumbent policy persistence.¹³ Let Σ_J be the set of Markov strategies for party J . Henceforth, the term strategy always refers to Markov strategies. With slight abuse of notation, the *state path* $\{(I^i, x^i)\}_{i=0}^\infty$ induced by profile (σ_L, σ_R) starting from (I^0, x^0) is defined recursively by

$$\begin{aligned} (I^0, x^0) &= (I^0, x^0), \\ (I^i, x^i) &= \tau((I^{i-1}, x^{i-1}), \sigma_{-I^{i-1}}(I^{i-1}, x^{i-1})). \end{aligned}$$

The *policy path* $\{x^i\}_{i=0}^\infty$ induced by (σ_L, σ_R) starting from (I^0, x^0) is the policy sequence of the corresponding state path. Discounted payoffs to the current opposition and incumbent parties from policy path $\{x^i\}_{i=0}^\infty$ induced by (σ_L, σ_R) starting from (I^0, x^0) are given by

$$\begin{aligned} V_{-I^0}(\sigma_{-I^0}, \sigma_{I^0}; (I^0, x^0)) &= \sum_{i=1}^{\infty} \delta_{-I^0}^{i-1} u_{-I^0}(x^i), \\ V_{I^0}(\sigma_{I^0}, \sigma_{-I^0}; (I^0, x^0)) &= \sum_{i=1}^{\infty} \delta_{I^0}^{i-1} u_{I^0}(x^i). \end{aligned}$$

Definition 2. A *Markov perfect equilibrium* is a strategy profile (σ_L, σ_R) such that, for each state (I, x) ,

$$\sigma_{-I}(I, x) \in \arg \max_{\sigma'_{-I} \in \Sigma_{-I}} V_{-I}(\sigma'_{-I}, \sigma_I; (I, x)).$$

Henceforth, the term equilibrium always refers to Markov perfect equilibrium. The one-deviation property allows the following characterisation of Markov perfect equilibria: (σ_L, σ_R) is an equilibrium if and only if, for each state (I, x) ,

$$\begin{aligned} \sigma_{-I}(I, x) \in \arg \max_{z \in X} \Big\{ & \mathbf{1}_{z \in W(I, x)} [u_{-I}(z) + \delta_{-I} V_{-I}(\sigma_{-I}, \sigma_I; (-I, z))] \\ & + \mathbf{1}_{z \notin W(I, x)} [u_{-I}(x) + \delta_{-I} V_{-I}(\sigma_{-I}, \sigma_I; (I, x))] \Big\}. \end{aligned}$$

The restriction to Markov strategies limits implicit equilibrium coordination by conditioning strategies only on current states and not on entire playing histories. To shed light on parties' long-run interactions, it seems preferable to assume that challengers' behaviour depends

¹³There are other ways to define Markov strategies in my model. For example, a strategy could specify an action only to opposition parties. The above definition simplifies notation in later definitions.

on incumbents' policies only insofar as they affect available winning policies. Given that parties square off in elections that are years apart and often involve different politicians, strategies that differentiate between events that occurred even a few elections ago would have problematic interpretations.

3 Outcomes Without Incumbent Policy Persistence

My model's distinctive policy dynamics are due to incumbent policy persistence, which transfers competition from *within* to *across* elections and dampens the incentives that lead to median convergence in standard models. To illustrate this, consider instead the repeated game in which incumbent and opposition parties simultaneously commit to policies. This stage game is the standard model of electoral competition between policy-motivated parties. To make this game history independent, the tie-breaking rule cannot depend on the identity of the incumbent party, and any random and nondegenerate tie-breaking rule will do. As is well known, the unique Nash equilibrium of the stage game has each party commit to the median policy.¹⁴ Call the repeated simultaneous move game the *model without incumbent policy persistence*.

Proposition 1. *In the unique subgame perfect equilibrium of the model without incumbent policy persistence, parties commit to the median policy after all histories.*¹⁵

In the model without policy persistence, any party J can enforce the policy path $(M, M, \dots)$ after some history by committing to the median policy in all elections. Hence any subgame perfect equilibrium must yield party J at least the payoff $\frac{1}{1-\delta_J}u_J(M)$ following all histories. Since this payoff results from the repetition of the stage game's unique Nash equilibrium, this is in fact J 's lowest subgame perfect equilibrium payoff. Concavity of parties' payoffs ensures that when policies deviate from the median, one party loses more than the other gains, which implies that any equilibrium path that differs from the median after some history cannot simultaneously guarantee payoffs of $\frac{1}{1-\delta_J}u_J(M)$ to J and $\frac{1}{1-\delta_{-J}}u_{-J}(M)$ to $-J$.

Proposition 1 shows that all but one of the long-run equilibrium outcomes of the model with incumbent policy persistence fail to arise in the absence of this assumption. In the model with incumbent policy persistence, parties can enforce policy path $(M, M, \dots)$ *only* following histories in which they can commit to new policies. In an equilibrium in which parties alternate in office, this opportunity occurs every other period. An opposition party J needs to be guaranteed a payoff of at least $\frac{1}{1-\delta_J}u_J(M)$ when it contests an election, but this is not the case one period later when J is an incumbent. It is the asynchronicity of policy choices under incumbent policy persistence along with discounting which allow deviations from the median. While parties understand that they may be worse off relative to policy path $(M, M, \dots)$ as an incumbent, they are offered sufficient

¹⁴For a survey, see Osborne (1995).

¹⁵Proofs of all results are contained in the Appendix.

payoff in their first term of office to balance this (discounted) loss. Policy paths exhibit a form of dynamic inconsistency for incumbents since if they could, they would prefer to commit to the median policy.

4 Outcomes With Incumbent Policy Persistence

This section presents my results for model with incumbent policy persistence.

4.1 Myopic Strategies

The myopic models of Wittman (1977) and Kramer (1977) provide two different benchmark strategy profiles. Both fail to constitute Markov perfect equilibria when parties are forward-looking but they illustrate, in special cases, the main incentives generated by the model.

Given state (R, r) with $r > M$, the unique optimal choice in the stage game for opposition party L is to commit to $\max\{0, 2M - r\}$, its most extreme winning policy. If $r \leq M$, committing to policy r , committing to a losing policy and choosing not to contest the election are all optimal. These are the optimal actions of Wittman's (1977) myopic policy-motivated parties. Define myopically optimal strategy for party L , σ_L^{my} , as

$$\sigma_L^{my}(R, r) = \begin{cases} \max\{0, 2M - r\} & \text{if } r \geq M, \\ Out & \text{if } r < M. \end{cases}$$

Party R 's myopically optimal strategy, σ_R^{my} , is defined symmetrically. In what follows, party J will be said to commit to moderate policies in state (I, x) if $\sigma_J(I, x) \in W(I, x)$ and $\sigma_J(I, x)$ is not J 's myopically optimal action at (I, x) .

In general, myopic strategies will fail to constitute an equilibrium since when faced with an incumbent whose policy is sufficiently extreme the opposition party finds it optimal to sacrifice present payoffs and commit to a moderate policy. This guarantees that in future elections, the party faces more moderate opponents. In fact, Proposition 3 establishes the precise (yet restrictive) condition under which myopic behaviour may be dynamically optimal.

On the other hand, a naive extension of median convergence results to my model has opposition parties commit to the median policy in all states. These are the optimal actions of Kramer's (1977) myopic office-motivated and vote-maximizing parties and they correspond to the unique subgame perfect equilibrium strategies of the model without incumbent policy persistence from Section 3. These strategies fail to constitute an equilibrium since an opposition party expecting a moderate opponent when it seeks reelection has no incentive to win the current election with a moderate policy. The sole cost of winning an election with non-moderate policies is the extremism it may generate in an opponent's policies in future elections, which is not a concern if an opponent never campaigns with extreme policies.

4.2 Equilibrium Policy Dynamics: Alternation

Long-run policy outcomes are defined, naturally, as limit points of sequences of policies induced by equilibrium dynamics from some initial state.

Definition 3. Policy y is a *long-run policy outcome under equilibrium* (σ_L, σ_R) starting from (I, x) if y is a limit point of the policy path induced by (σ_L, σ_R) starting from (I, x) .

Policy y is an *alternating long-run policy outcome under equilibrium* (σ_L, σ_R) starting from (I, x) if there exists policy $z \neq y$ such that both y and z are limit points of the policy path induced by (σ_L, σ_R) starting from (I, x) .

A policy that is a long-run policy outcome under some equilibrium starting from some state will be called simply a *long-run policy outcome*. Alternating long-run policy outcomes capture dynamics in which alternation in power occurs infinitely often, since under incumbent policy persistence policy change occurs only when an incumbent is defeated. Proposition 2 characterises equilibrium dynamics along with the properties of their limit points.

Proposition 2. Consider some equilibrium (σ_L, σ_R) and some state (I, x) along with the policy path $\{y^i\}$ induced by (σ_L, σ_R) starting from (I, x) . Suppose that $(I, x) = (R, r)$.

- i. If $r \leq M$, then $y^i = r$ for all i .
- ii. If $r > M$, then incumbents are always defeated on the equilibrium path and $\{y^i\}$ has a pair of limit points $(\hat{\ell}, 2M - \hat{\ell})$ for some $\hat{\ell} \leq M$. Furthermore, $\sigma_L(R, 2M - \hat{\ell}) = \hat{\ell}$ and $\sigma_R(L, \hat{\ell}) = 2M - \hat{\ell}$.

The case of $(I, x) = (L, \ell)$ is symmetric.

While item i of Proposition 2 shows that all policies are long-run policy outcomes, item ii ensures that all nontrivial equilibrium dynamics entail alternation in power and convergence to symmetric pairs of policies of the form $(\ell, 2M - \ell)$ for some $\ell < M$. Figure 2 illustrates this result, depicting a policy path induced by some strategy profile from state (R, r^0) with $r^0 > M$. The policy path alternates around the median and converges to an alternation around the pair (ℓ, r) , which is an alternating long-run policy outcome from (R, r^0) if the strategy profile is an equilibrium. The pair of long-run policies (ℓ, r) need not be reached by the policy path.

In any equilibrium, party L will stay out, or commit to some losing policy, whenever (R, r) is such that $r < M$, that is, when party R is on the left of the political spectrum. The policy path most favourable to L that can be sustained in any equilibrium from such a state is $(r, r, r, \dots)$, which L can attain by staying out and trapping dynamics at the initial policy forever. In turn party R , in states (L, ℓ) with $\ell < M$, will never select a winning policy strictly to the left of M . On any equilibrium path starting from a state in which the incumbent party champions a policy on its side of the median, power changes hands each election and opposition parties commit to

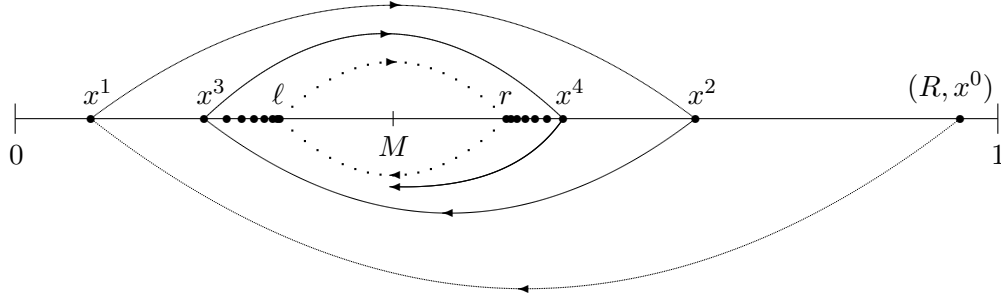


Figure 2: Example of Equilibrium Policy Dynamics.

winning policies on their side of the median. Such policy paths have at most two limit points since the sequences of each party's winning policies are monotone. This limiting outcome must be a pair of policies symmetric about the median (including the median itself). Item ii of Proposition 2 shows that limits of alternating equilibrium dynamics are absorbing; if the dynamics start at one of the limiting policies, they stay there.

The proof of Proposition 2, along with the results to follow, depend only on properties of parties' preferences over alternating long-run policy outcomes. Parties' payoffs to sequences of alternating symmetric policies vary according to the initial policy. To clarify this, define the functions $\{\Psi_L^\theta : [0, M] \rightarrow \mathbf{R}\}_{\theta \in \{+, -\}}$ for party L as

$$\begin{aligned}\Psi_L^+(\ell) &= u_L(\ell) + \delta_L u_L(2M - \ell) \text{ and} \\ \Psi_L^-(\ell) &= u_L(2M - \ell) + \delta_L u_L(\ell).\end{aligned}$$

The discounted sum $\frac{1}{1-\delta_L^2}\Psi_L^+(\ell)$ is party L 's payoff from the alternation of policies $(\ell, 2M - \ell)$ starting from ℓ , while $\frac{1}{1-\delta_L^2}\Psi_L^-(\ell)$ is its payoff to the same alternation when starting from $2M - \ell$. Functions $\{\Psi_R^\theta : [0, M] \rightarrow \mathbf{R}\}_{\theta \in \{+, -\}}$ for party R are defined symmetrically. Strict concavity of parties' utility functions yields a natural preference order over symmetric alternations. The following lemma, an immediate consequence of concavity, is stated here for the sake of clarity.

Lemma 1. *There exist uniquely defined policies ℓ^* and r^* such that*

$$\begin{aligned}\ell^* &= \arg \max_{\ell \in [0, M]} \Psi_L^+(\ell) \in [0, M) \text{ and} \\ r^* &= 2M - \arg \max_{\ell \in [0, M]} \Psi_R^+(\ell) \in (M, 2M].\end{aligned}$$

Ψ_L^- (Ψ_R^-) is strictly increasing (decreasing) and both Ψ_J^+ and Ψ_J^- are strictly concave for $J \in \{L, R\}$.

Given $\ell \in [0, M)$, the concavity of u_L ensures that the cost to L of a moderate move away from ℓ is dominated by the benefit of a moderate move away from $2M - \ell$. That Ψ_L^+ is single-peaked around $\ell^* < M$ while Ψ_L^- is strictly increasing follow from discounting. When the payoff to L from alternating pairs are evaluated from L 's policy, a shift to a more moderate alternation ensures that party L suffers the full loss to moderation in its own policy, while the larger benefit of R 's moderation is discounted. Hence, more moderate alternations are preferred until $(\ell^*, 2M - \ell^*)$, which is L 's favoured alternation when starting from its own policy. For any $\delta_L < 1$, ℓ^* is bounded away from the median as $\lim_{\ell \nearrow M} u_L(\ell) = \lim_{\ell \nearrow M} u_L(2M - \ell)$. Policies ℓ^* and r^* are central in the characterisation of alternating long-run policy outcomes in the next section. Meanwhile, when the payoffs to L from alternations are evaluated from R 's policy, L always prefers more moderate alternations, and in particular, its favoured alternation is that around M . This follows since L 's loss from moderating its own policy, smaller than its gain from R moderating its own policy, is discounted.

When $\ell^* > 0$, it is uniquely determined by $\frac{u'_L(\ell^*)}{u'_L(2M - \ell^*)} = \delta_L$ and increasing in δ_L since u_L is concave. As party L becomes less short-sighted, the cost of R 's future policies increases and its preferred alternation comes closer to the median, converging to M as δ_L converges to 1. Similarly, ℓ^* is increasing in L 's disutility for policies away from its ideal point. Consider v_L obtained from u_L by applying some increasing concave transformation. Then for any $\ell \in (0, M)$, $\frac{v'_L(\ell)}{v'_L(2M - \ell)} < \frac{u'_L(\ell)}{u'_L(2M - \ell)}$, and hence ℓ^* approaches M as parties' utilities become more concave.

4.3 Alternating Long-Run Policy Outcomes: Bounded Extremism

Proposition 3 characterises alternating long-run policy outcomes and shows that they display *bounded extremism*.

Proposition 3. *Policies $(\ell, 2M - \ell)$ for some $\ell \leq M$ are alternating long-run policy outcomes if and only if $\ell \in [\max\{\ell^*, 2M - r^*\}, M]$.*

Figure 3, illustrates Proposition 3 when $\ell^* \geq 2M - r^*$. The dotted section of the policy space indicates the set of long-run policy outcomes. All symmetric policy pairs more moderate than $(\ell^*, 2M - \ell^*)$, such as (ℓ, r) and (ℓ', r') , are alternating long-run policy outcomes, with $(\ell^*, 2M - \ell^*)$ the most extreme such outcome.

The bound on long-run extremism established by Proposition 3 is due only to parties' incentives to manage the sets of winning policies their opponents inherit from their time in office. It follows for the same reason, given in Section 4.1, that the myopically optimal strategy profile fails to constitute an equilibrium: facing a sufficiently extreme alternation (in the long-run), some party will prefer to rein in future opponents' policies by committing to more moderation.

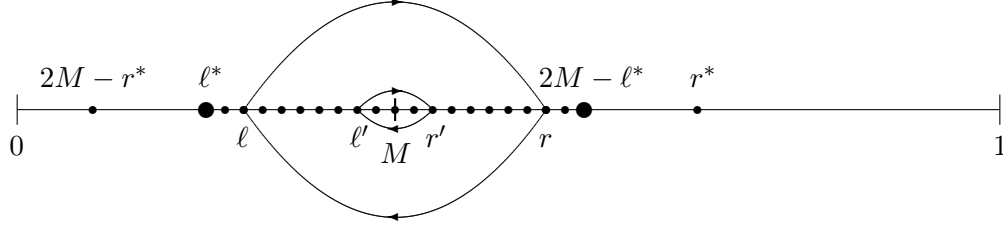


Figure 3: Set of Alternating Long-run Policy Outcomes.

Proposition 3 makes precise how long-run deviations from the median are driven by parties' myopic preference for their own policies. If parties are close to myopic, a wide range of alternating outcomes are supported in equilibrium. However, if only one party is arbitrarily far-sighted, policies are arbitrarily close to the median in the long-run since even if long-run outcomes are symmetric, which ones are achievable is determined by the party that prefers more moderate policy alternations.

There are two steps to proof of Proposition 3. Lemma 4 establishes the existence of the bound on extremism $\max\{\ell^*, 2M - r^*\}$. Its proof hinges on a useful lower bound party L 's equilibrium payoff: any equilibrium path following a commitment to some winning policy ℓ yields a payoff of at least $\frac{1}{1-\delta_L^2} \Psi_L^+(\ell)$. To see this, consider a strategy for opposition party L which sets policy ℓ in the current election and responds myopically to all of R 's subsequent policies. The payoff to L from this strategy is $u_L(\ell)$ in this election, along with a sequence of payoffs $\{\Psi_L^-(r^i)\}$ in the subsequent pairs of elections, for some sequence of policies $\{r^i\}$ such that $\ell \leq 2M - r^i$ for all i . By Lemma 1, each payoff in this sequence is at least $\Psi_L^-(\ell)$ and hence the payoff to selecting winning policy ℓ must be at least $\frac{1}{1-\delta_L^2} \Psi_L^+(\ell)$. Hence, if $(\ell, 2M - \ell)$ were an alternating long-run policy outcome with $\ell < \ell^*$, party L could win the election in state $(R, 2M - \ell)$ by committing to policy ℓ^* and guarantee itself a payoff of at least $\frac{1}{1-\delta_L^2} \Psi_L^+(\ell^*)$, its preferred alternation. However, L 's equilibrium payoff in state $(R, 2M - \ell)$ is $\frac{1}{1-\delta_L^2} \Psi_L^+(\ell)$, yielding the required contradiction.

Lemma 6 completes the proof of Proposition 3 by showing that the bound on long-run extremism is tight. This consists of constructing an equilibrium under which all alternations $(\ell, 2M - \ell)$, with $\ell \geq \max\{\ell^*, 2M - r^*\}$ are alternating long-run policy outcomes. Consider the strategy $\sigma_L^{\ell^*}$

such that

$$\sigma_L^{\ell^*}(R, r) = \begin{cases} \ell^* & \text{if } r \in [2M - \ell^*, 1], \\ 2M - r & \text{if } r \in [M, 2M - \ell^*), \\ Out & \text{if } r \in [0, M). \end{cases}$$

Lemma 6 shows that if $\ell^* \geq 2M - r^*$, $(\sigma_L^{\ell^*}, \sigma_R^{my})$ is an equilibrium. If $\ell^* < 2M - r^*$, strategy $\sigma_R^{r^*}$ can be defined by reversing then the roles of the two parties and it can be shown that $(\sigma_L^{my}, \sigma_R^{r^*})$ is an equilibrium. These equilibria provide the exact condition under which parties' behaviour in Wittman (1977) can be said to be dynamically rational: myopically optimal strategies form an equilibrium if and only if $\max\{\ell^*, 2M - r^*\} = 0$.

Figure 4, depicting the interval $[0, M]$, illustrates the equilibrium for the case of $\ell^* \geq 2M - r^*$. The directed curve above (below) the interval from point ℓ represents the equilibrium action of party L (R) in state $(R, 2M - \ell)$ ((L, ℓ)). In equilibrium, from any (L, ℓ) with $\ell < \ell^*$ or (R, r) with $r > 2M - \ell^*$ equilibrium policies settle on alternation $(\ell^*, 2M - \ell^*)$ in at most two elections. Note that policies $(\ell, 2M - \ell)$ with $\ell \geq \ell^*$ are long-run policy outcomes, but only starting from (L, ℓ) or $(R, 2M - \ell)$.

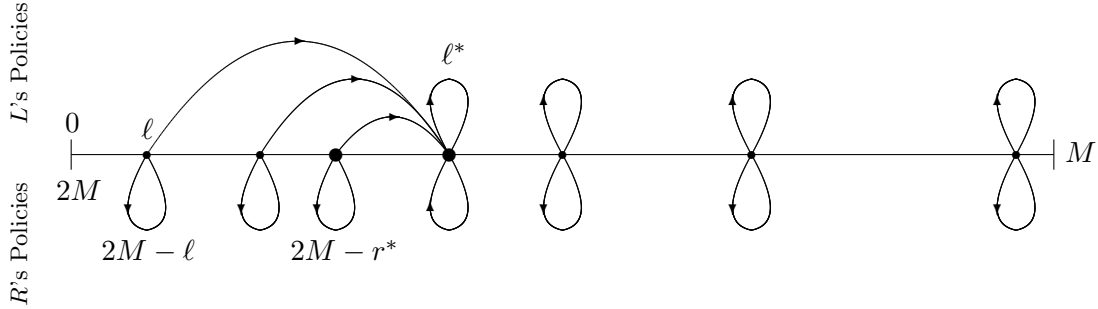


Figure 4: Policy Dynamics of MPE $(\sigma_L^{\ell^*}, \sigma_R^{my})$ when $\ell^* \geq 2M - r^*$.

In equilibrium, responding myopically in all states is dynamically rational for party R . When the state is (L, ℓ) for some $\ell < \ell^*$, party R knows party L will commit to ℓ^* in the next election against any winning policy $r \in [2M - \ell^*, 2M - \ell]$ it champions in the current election. Hence committing to myopic policy $2M - \ell$ is optimal. Meanwhile if the state is (L, ℓ) for some $\ell \geq \ell^*$, party R knows that championing any policy $r \in [M, 2M - \ell]$ in the current election induces an alternation around $(2M - r, r)$. Since $2M - \ell \leq 2M - \ell^* \leq r^*$, R 's preferred alternation among those it can enforce is $(2M - \ell, \ell)$.

In moderate states (L, ℓ) for some $\ell \geq \ell^*$ and (R, r) for some $r \leq 2M - \ell^*$, both parties respond myopically. For these states, their preferences over alternations coincide; both prefer more extreme alternations when evaluated from their own policy. Parties' preferences over alternations also coincide in extreme states (L, ℓ) for some $\ell < 2M - r^*$ and (R, r) for some $r > r^*$. In these states both parties prefer more moderate alternations when evaluated from their own policy. However, having both parties committing to moderate policies in these states cannot be an equilibrium and some party must bear the 'responsibility' of bringing policy dynamics towards more moderate alternations. For intermediate states (L, ℓ) for some $\ell \in [2M - r^*, \ell^*)$ and (R, r) for some $r \in (2M - \ell^*, r^*]$, parties' preferences over alternations diverge and party L , which prefers more moderate pairs, ensures that policy paths converge.

4.4 Robust Long-run Policy Outcomes: Bounded Moderation

A long-run policy outcome y is the limit of equilibrium policy dynamics given *some* initial state. In particular, 'steady state' outcome y need not be dynamically stable in the following sense: given an initial state with policy more extreme than either ℓ or r , equilibrium policy dynamics need not converge to (ℓ, r) . The following definition addresses this issue.

Definition 4. Policy y is a *robust long-run policy outcome* if it is a long-run policy outcome under some equilibrium (σ_L, σ_R) starting from some state (I, x) such that x is not a limit point of the policy path induced by (σ_L, σ_R) from (I, x) .

Robustness is a weak requirement of dynamic stability as it necessitates only the existence of a single policy x that lies on an equilibrium path that has y as a limit point. However, only alternating long-run policy outcomes and the median outcome M can be robust, as all other long-run policy outcomes arise only from a single state.

The main result of this section is a characterisation of the set of robust long-run policy outcomes under equilibria in *consistent strategies*.¹⁶ These strategies require that parties do not condition on the exact policy of the incumbent when choosing policies in the interior of their set of winning policies. A party which commits to a moderate policy is unconstrained by the incumbent's policy, and hence facing a slightly more moderate incumbent should not lead it to change its policy choice. Figure 5 illustrates this requirement. Suppose that party L chooses winning policy $\ell > 2M - r$ from set of winning policies $[2M - r, r]$ for some $r > M$. Consistency requires that Party L choose the same policy from a set of winning policies $[2M - r', r']$ for some $2M - r' < \ell$. The choice of any other policy from the smaller set of winning policies could be justified by equilibrium considerations, but not by any fundamental political constraints.

Definition 5. A *consistent Markov strategy* is a Markov strategy such that

¹⁶A class to which, notably, the equilibria of Section 3 belong.

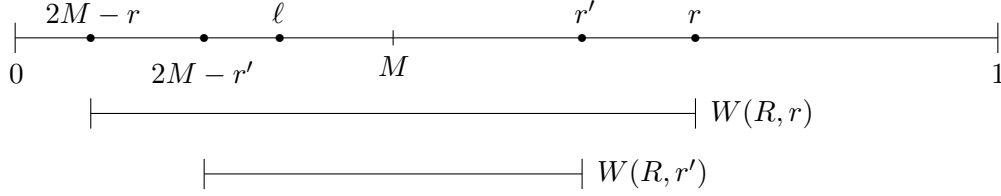


Figure 5: Consistent Markov Strategies.

- i. For $x, x' \in X$ such that $W(I, x') \subset W(I, x)$, suppose $\sigma_{-I}(I, x) \in W(I, x)$. If $\sigma_{-I}(I, x) \in W(I, x')$ then $\sigma_{-I}(I, x') = \sigma_{-I}(I, x)$.
- ii. For $x \in X$, say $\sigma_{-I}(I, x) \notin W(I, x)$. Then $\sigma_{-I}(I, x) = Out$ and $\sigma_{-I}(I, x') = Out$ for all $x' \in X$ such that $u_{-I}(x') > u_{-I}(x)$.

A *consistent Markov perfect equilibrium* is a Markov perfect equilibrium in consistent Markov strategies.

Requirement i expresses the restriction explained above formally. Requirement ii normalises losing policy choices to non-participation and requires that if party L selects *Out* in state (R, r) , then it does so for all states (R, r') with $r' < r$.

Consistency restricts equilibrium coordination and in this sense is a strengthening of the requirement that Markov strategies depend solely on payoff relevant information. It goes beyond payoff relevance since the histories that are considered relevant to a party's decision are defined relative to its strategy. In the example above, party L *reveals*, through its choice of ℓ in state (R, r) , that states (R, x) for $x \in [2M - \ell, r)$ are of no intrinsic importance to its decisions. Also note that by the one-deviation property, the restriction to consistent equilibria does not affect the deviations available to parties, but only equilibrium strategies. That is, a consistent equilibrium is an equilibrium. Proposition 4 characterises robust long-run outcomes under consistent strategies and shows that they display *bounded moderation*.

Proposition 4. *There exists $\ell^{**} \in (\max\{\ell^*, 2M - r^*\}, M)$ such that policies $(\ell, 2M - \ell)$ for some $\ell \leq M$ are robust long-run policy outcomes under equilibria in consistent strategies if and only if $\ell \in [\max\{\ell^*, 2M - r^*\}, \ell^{**}]$.*

Figure 6 illustrates Proposition 4. The dotted line indicates the set of alternating long-run policy outcomes, while the dashed line indicates the subset of these pairs that are robust under

consistent equilibria in consistent strategies. For example, policy pair (ℓ, r) is robust, while policy pair (ℓ', r') , more moderate than $(\ell^{**}, 2M - \ell^{**})$, is not.

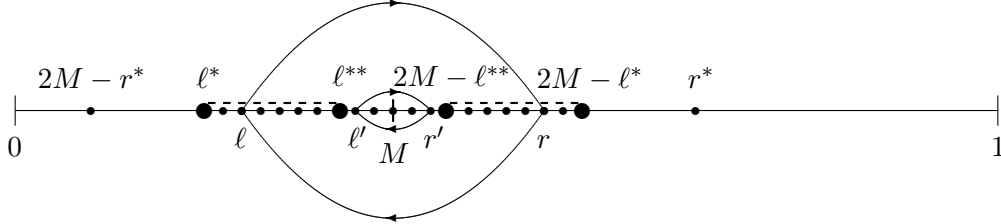


Figure 6: Set of Robust Long-run Policy Outcomes under Consistent Equilibria.

The model admits median politics as a long-run policy outcome only if the initial incumbent party champions the median. This does not contradict the results of Section 4; centripetal forces are present and policy paths tend to converge *toward* the median. However, policies do not converge *to* the median. Parties remain differentiated in the long-run and settle into clearly defined party identities. On equilibrium convergence paths a party's commitment to a moderate policy must be reciprocated in future elections by its opponent. When converging to sufficiently moderate policy alternations, parties' value their opponents' (discounted) moderate moves so little that they are unwilling to commit to policies moderate enough to sustain convergence.

When studying the convergence outcomes of the model, it is convenient to focus on the symmetric images of party R 's policies with respect to the median, mapping converging dynamics into a single increasing sequence of policies. The *convergence path* $\{y^i\}$ to policy $\hat{\ell} \in (0, M]$ under equilibrium (σ_L, σ_R) starting from (I^0, x^0) is a sequence such that

- i. If $(I^0, x^0) = (R, r)$ for some $r > 2M - \hat{\ell}$, then $y^0 = 2M - x^0$, $y^i = x^i$ for i odd and $y^i = 2M - x^i$ for i even, where $\{x^i\}$ is the sequence of policies induced by (σ_L, σ_R) starting from (I^0, x^0) .
- ii. If $(I^0, x^0) = (L, \ell)$ for some $\ell < \hat{\ell}$, then $y^0 = \ell$, $y^i = x^i$ for i even and $y^i = 2M - x^i$ for i odd.
- iii. $\{y^i\} \rightarrow \hat{\ell}$.

Consistent strategies allow simple characterisations of parties' policy choices and payoffs on equilibrium convergence paths. Lemma 7 characterises strategies along convergence paths in

consistent strategies and is illustrated in Figure 7, showing a section of some convergence path $\{y^i\}$ initiated by party R committing to policy $2M - y^i$, to which L responds by moderating to y^{i+1} . By the definition of consistent strategies, $\sigma_L(R, r) = y^{i+1}$ for all $r \in (2M - y^i, 2M - y^{i+1}]$, that is, L moderates to y^{i+1} when facing an incumbent R championing a policy more moderate than $2M - y^i$. Furthermore, consistency implies that $\sigma_R(L, \ell') = 2M - \ell'$ for all $\ell' \in [y^i, y^{i+1})$, that is, R responds myopically whenever L stops short of moderating to y^{i+1} .¹⁷

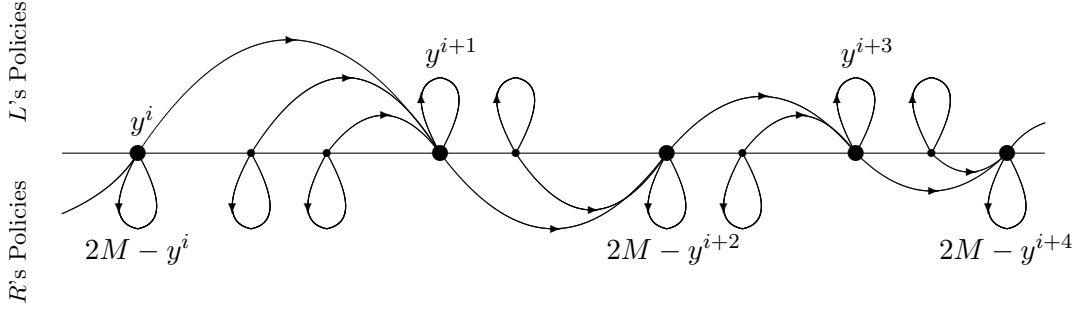


Figure 7: Convergence Paths under Consistent Equilibria.

As noted in Section 4.1, it is optimal to respond myopically to an opponent that always selects the median policy. Figure 8 shows that consistent equilibria display this behaviour locally. That is, consistent equilibrium convergence paths define alternating sets of policies in which a locally myopic party meets a locally moderate party. Parties stake out non-negotiable ‘core’ issues and their opponents compromise on the corresponding policies on the other side of the median. The location of parties’ core issues may seem idiosyncratic since they compromise over neighbouring policies. However, core issues are not due to parties’ preferences for particular policies but arise as tool to sustain policy convergence.

Section 4.3 noted that $\frac{1}{1-\delta_L^2} \Psi_L^+(y^i)$ is a lower bound on L ’s payoff in state $(R, 2M - y^i)$. Lemma 8 shows that if $2M - y^i$ lies on a consistent equilibrium convergence path then this payoff is also an upper bound. That is, L ’s payoff at $(R, 2M - y^i)$ is computed ‘as though’ equilibrium dynamics were absorbed by the symmetric pair of policies $(y^i, 2M - y^i)$. In consistent equilibria, opposition parties are indifferent between sustaining convergence and enforcing a symmetric alternation at the current policy. However, since in state $(R, 2M - y^{i+2})$ party L receives the

¹⁷Note that if policies $(\hat{\ell}, 2M - \hat{\ell})$ are robust long-run policy outcomes under consistent equilibria and $y^i < \hat{\ell}$ is on a convergence path to $\hat{\ell}$, then all $\ell \in (y^i, \hat{\ell})$ are also on a convergence path to $\hat{\ell}$. In this sense, convergence outcomes under consistent equilibria can be said to be ‘strongly’ robust since convergence to $\hat{\ell}$ occurs from all more extreme states.

payoff to an alternation at $(y^{i+2}, 2M - y^{i+2})$, its payoff to being in opposition on convergence paths to policies more moderate than ℓ^* is strictly decreasing. After each spell in opposition, parties regret their previous moderate policies. L 's payoff in state $(R, 2M - y^i)$ satisfies

$$\Psi_L^+(y^i) - \Psi_L^+(y^{i+1}) = \delta_L[\Psi_L^-(y^{i+2}) - \Psi_L^-(y^{i+1})]. \quad (1)$$

The left-hand side of (1) is the cost (computed in payoffs to alternations starting from L 's policy) of choosing moderate policy y^{i+1} while the right-hand side is the (discounted) benefit (computed in payoffs to alternations starting from R 's policy) of party R 's subsequent moderation to $2M - y^{i+2}$. These costs and benefits are balanced by the choice of y^{i+1} . On convergence paths to policies more moderate than ℓ^* , moderation is costly and party L commits to y^{i+1} only because it expects R to continue to sustain convergence in the next election. That is, (1) cannot be satisfied for $y^{i+2} = y^{i+1}$ unless $y^{i+1} = y^i$. Moderation is self-reinforcing: if parties anticipate an end to convergence in the future current incentives to choose moderate policies unravel.¹⁸ Equation (1) also explains why party L is willing to sustain convergence paths to alternations more moderate than $(\ell^*, 2M - \ell^*)$, that is, why $\ell^{**} > \ell^*$. Around ℓ^* , cost of moving to a more moderate alternation is of second-order importance, while the benefit of R 's moderation is of first-order importance. Around ℓ^* , L is willing to bear much of the cost of sustaining convergence.

The recursive relationship in (1), along with the corresponding relationship for party R , allow the derivation of the bound on moderation $\ell^{**} \in (\ell^*, M)$. Fix one round of moderation from $(R, 2M - y^i)$ as the moves, first by L , then by R , that take the state to $(R, 2M - y^{i+2})$. Then (1) describes the share of the total moderation $y^{i+2} - y^i$ that L is willing to undertake. The bound ℓ^{**} is the most moderate policy for which the parties 'supply' of moderation is consistent with convergence in the limit as $y^{i+2} \rightarrow y^i$. Convergence to moderate policies fails as the shares of any given round of moderation that parties are willing to undertake become too small. To see this, consider the polar case of convergence to the median. As some convergence path approaches M , moderate moves of similar sizes by parties L or R have similar effects on L payoffs, yet the gain from R 's moderation is discounted. Since the same observation holds for R , both parties require their opponents to moderate more than they do, which contradicts convergence.

As in section 4.3, the bound ℓ^{**} is shown to be tight through the construction of equilibria. Contrary to that section, where a single equilibrium yields all alternating long-run policy outcomes, an equilibrium under which policies $(\hat{\ell}, 2M - \hat{\ell})$ are robust long-run policy outcomes is constructed for each $\hat{\ell} \in (\max\{\ell^*, 2M - r^*\}, \ell^{**}]$. For such $\hat{\ell}$, given a policy path $\{y^i\}$ such that $y^0 = \ell^*$, $\{y^i\} \rightarrow \hat{\ell}$ and satisfying (1), Lemma 12 provides the equilibrium strategies and verifies their optimality. As in equilibrium $(\sigma_L^{\ell^*}, \sigma_R^{my})$, policies from any state more extreme than ℓ^* move rapidly to ℓ^* , and from there a convergence path ensures they approach $\hat{\ell}$.

¹⁸In fact, this holds for all equilibria, not just those in consistent strategies. For the same reasons as above, but without relying on payoff condition (1), it can be shown that only the most extreme alternating outcomes, $(\max\{\ell^*, 2M - r^*\}, 2M - \max\{\ell^*, 2M - r^*\})$, are ever reached from a more extreme state in a finite number of elections in any equilibrium.

Lemma 13 shows that a sequence as specified above exists. The proof consists in iterating the recursive relationship in (1) forward from $y^0 = \ell^*$ through the choice of y^1 and establishing the conditions under which this operation defines a converging policy path. Given any $\hat{\ell} \in (\ell^*, \ell^{**}]$, some policy $y^1 > \ell^*$ can be found such that $\{y^i\} \rightarrow \hat{\ell}$. From above, when $\hat{\ell} < \ell^{**}$ the share of moderation around $\hat{\ell}$ that parties are willing to undertake exceeds the amount of moderation that needs to be allocated to sustain convergence. The result hinges on the concavity of Ψ_L^+ and Ψ_L^- , as this ensures that parties become less willing to compromise as policies get closer to the median and hence the share of moderations that parties are willing to undertake at all ℓ' with $\ell' < \ell < \ell^{**}$ are larger than those they are willing to undertake at ℓ .

5 Discussion and Extensions

5.1 Forward-looking Voters

Myopic voting guarantees that all future governments are at least as moderate as the current incumbent. However, forward-looking voters may choose to elect opposition parties with more extreme platforms than incumbents if this can guarantee preferred continuations. On an alternating equilibrium path, the median voter has no incentive to support the incumbent since by voting against a (weakly) more moderate opposition, it is worse off in this election and faces the same choice in the next election. More generally, the median voter approves of converging equilibrium outcomes, and hence voting for opposition parties that propose moderate policies is in its interest. However, equilibria with myopic voters do not persist as equilibria of a game in which voters are forward-looking. The previous discussion suggests that the issues with myopic voting arise off the equilibrium convergence path.

Consider the extension of the model in which voters are forward-looking. I restrict attention to equilibria in which the median voter is decisive, and consider a single representative median voter with utility function u_M and discount factor δ_M . A strategy for the voter is $\sigma_M : (\{L, R\} \times X) \times (X \times \{Out\}) \rightarrow \{0, 1\}$, where $\sigma_M((I, x), z) = 0$ if and only if the median voter supports incumbent I with policy x in an election opposing it to $-I$ with policy z . Assume that the median voter never abstains so that in particular $\sigma_M((I, x), Out) = 0$ for all (I, x) . Denote the set of strategies for M as Σ_M . As in Section 2, a profile of strategies $(\sigma_L, \sigma_R, \sigma_M)$ along with state (I^0, x^0) determines discounted payoff $V_J(\sigma_L, \sigma_R, \sigma_M; (I^0, x^0))$ for player $J \in \{L, R, M\}$.

Definition 6. A *Markov perfect equilibrium with forward-looking voters* is a strategy profile $(\sigma_L, \sigma_R, \sigma_M)$ such that for each state (I, x)

$$\sigma_{-I}(I, x) \in \arg \max_{\sigma'_{-I} \in \Sigma_{-I}} V_{-I}(\sigma'_{-I}, \sigma_I, \sigma_M; (I, x)),$$

and for any policy z ,

$$\sigma_M((I, x), z) \in \arg \max_{\sigma'_M \in \Sigma_M} V_M(\sigma_L, \sigma_R, \sigma'_M; (I, x)).$$

Equilibria in consistent strategies for parties are not equilibria with forward-looking voters. To see this, consider a consistent equilibrium convergence path $\{y^i\}$, a policy y^i such that $\sigma_L(R, 2M - y^i) = y^{i+1}$, a state (L, y') for some $y' \in (y^i, y^{i+1})$ and a deviation by R to $2M - y' + \epsilon$ for some $\epsilon < y' - y^i$. The median voter's myopic strategy, σ_M^{my} , calls for a vote against R . If it does so, its payoff V_M^L is given by

$$V_M^L = u_M(y') + \delta_M u_M(2M - y') + \delta_M^2 V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (R, 2M - y')).$$

If instead the median voter votes for R , its payoff V_M^R is given by

$$V_M^R = u_M(2M - y' + \epsilon) + \delta_M V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (R, 2M - y')).$$

By symmetry of u_M ,

$$\lim_{\epsilon \rightarrow 0} (V_M^L - V_M^R) = \delta_M(1 - \delta_M) \left[\frac{1}{1 - \delta_M} u_M(2M - y^i) - V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (R, 2M - y^i)) \right] < 0,$$

since the equilibrium path following $(R, 2M - y^i)$ consists of a converging policy path all strictly closer to the median than $2M - y^i$. In the sections of the convergence path in which party R responds myopically, the median voter finds it costly to punish extreme deviations by party R . To do so, it must vote for the incumbent party L and keep it in power for another term, but this simply delays R 's victory by one period. By voting against party R , the median voter delays the resumption of convergence by two elections. Voting for deviating party R in this election lets a more moderate party L gain office in the next election.

Proposition 5 shows that given any alternating consistent equilibrium convergence path, it is possible to construct voter and party strategies that enforce this path in an equilibrium with forward-looking voters. Hence, all the equilibrium outcomes of this paper are robust to forward-looking voters.

Proposition 5. *Consider consistent equilibrium (σ_L, σ_R) in the game with myopic voters. Consider state (I, x^0) such that $I = L$ and $x^0 \leq M$ or $I = R$ and $x^0 \geq M$, along with policy path $\{y^i\}$ induced from (I, x^0) by (σ_L, σ_R) . Then there exist an equilibrium with forward-looking voters $(\sigma'_L, \sigma'_R, \sigma_M)$ such that the policy path $\{y^i\}$ induced from (I, x^0) by $(\sigma'_L, \sigma'_R, \sigma_M)$ is such that $y^i = y^i$ for all $i \geq 2$.*

In the equilibrium $(\sigma'_L, \sigma'_R, \sigma_M)$, the median voter sometimes votes against myopically preferred policies. In particular, it does so whenever the problem with consistent equilibrium convergence paths identified above arises. In the sections of the convergence path in which party R responds myopically, the median voter supports more extreme policies to ensure a quicker defeat by R at the hands of L and a resumption of convergence.

The equilibrium strategies are illustrated in Figure 8. Consider policy y^i such that $\sigma_L(R, 2M - y^i) = y^{i+1}$. In the Appendix, for $\ell \in [y^i, y^{i+1})$ I define function $z^{i+1}(\ell) \in [y^i, \ell)$ such that $2M - z^{i+1}(\ell)$ is the most extreme policy by R supported by the median voter against ℓ in state (L, ℓ) . Note that to ‘resume’ convergence, the median voter is never willing to support a policy by R that is more extreme than $2M - y^i$. Suppose, for example, that the median voter supported a proposal $r \in (2M - y^i, 2M - y^{i-1}]$. In the next election, L does not commit to a moderate policy (in fact it commits to $z^i(r) \in [y^{i-1}, r)$), and it takes two elections to return to state $(R, 2M - y^i)$. Against this, the median voter prefers to vote against R and wait two elections to arrive at the more moderate state (L, y^{i+1}) .

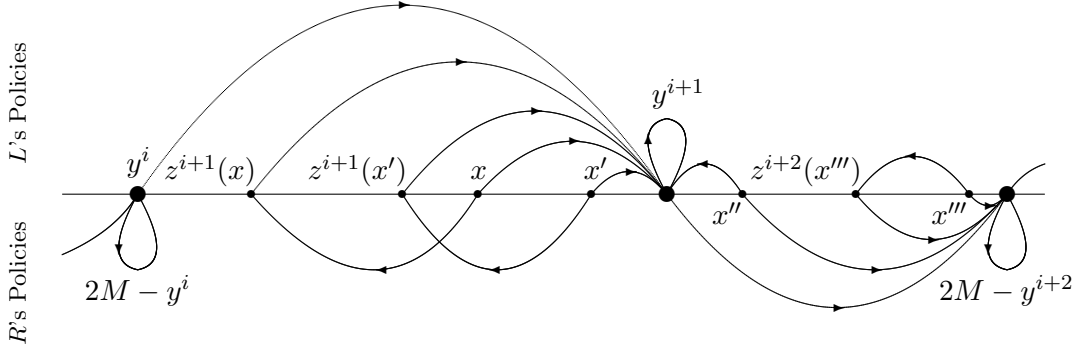


Figure 8: Policy Dynamics of Equilibria with Forward-looking Voters.

When voters are myopic, in states (L, ℓ) with $\ell > M$ party R does not participate in any elections and policy dynamics get trapped. When the median voter is forward-looking, it may vote in favour of policies by R that are more extreme than ℓ but lead to renewed convergence. Hence the result of Proposition 5 applies only to alternating convergence paths, and it needs to be shown that parties in alternating convergence paths have no incentive to commit to a policy on their opponent's side of the median solely to have convergence eventually resume from a more extreme initial state.

5.2 Legislative Bargaining

This section discusses in detail the relationship between my paper and dynamic legislative bargaining models, in particular that of Baron (1996), which features a single-dimensional policy space. There are three crucial differences between my paper and Baron's (1996). First, I do not assume that the median voter is represented by a party that shares its preferences over policies. Second, incumbent policy persistence generates a history-dependent proposer recognition rule. Finally, the equilibria I construct in Section 4.4 have the median voter strictly prefer to support opposition parties on the equilibrium path.

To view my model as a legislative bargaining model, suppose there is a large number of legislators, with M denoting the ideal policy of the median legislator. However, in contrast to Baron (1996), only two legislators can be recognised to propose policies; these are legislators L and R that have ideal policies 0 and 1. Assumed that they are recognised each period with equal probability. In the *legislative bargaining model*, a state (I, x) consists of the current proposer along with the status quo. A proposal strategy for party I is $\sigma_I : \{I\} \times X \rightarrow X$.¹⁹ Consider, as above, voting strategies σ_M for the median legislator, where now $\sigma_M((I, x), z) = 0$ if and only if M supports the status quo, and discounted payoffs $V_K(\sigma_L, \sigma_R, \sigma_M; (I, x))$ for player $K \in \{L, R, M\}$. An equilibrium of the legislative bargaining game is as in Definition 6, with the relevant reinterpretations. When legislators are forward-looking, the set of winning policies for I in state (I, x) is determined by σ_M . Define $W^{\sigma_M}(I, x)$ to be the set of winning policies for I in state (I, x) .

Definition 7. A *consistent Markov proposal strategy* is a Markov strategy such that

- i. For $x, x' \in X$ such that $W^{\sigma_M}(I, x') \subset W^{\sigma_M}(I, x)$, suppose $\sigma_I(I, x) \in W^{\sigma_M}(I, x)$. If $\sigma_I(I, x) \in W^{\sigma_M}(I, x')$ then $\sigma_I(I, x') = \sigma_I(I, x)$.
- ii. For $x \in X$, say $\sigma_I(I, x) \notin W^{\sigma_M}(I, x)$. Then $\sigma_I(I, x) = x$ and $\sigma_I(I, x') = x$ for all $x' \in X$ such that $u_I(x') > u_I(x)$.

A convergence path $\{y^i\}$ in the legislative bargaining game is as defined above. However, in this framework the description of a convergence path no longer corresponds to the realised equilibrium policy path. The Appendix shows that myopic voting is optimal for the median legislator when facing consistent strategies. Since it is without loss of generality to assume that in any equilibrium $y^{i+1} \in W^{\sigma_M}(L, y^{i+1})$, consistent Markov proposal strategies along with $\sigma_M = \sigma_M^{my}$ immediately imply that if $\sigma_L(L, 2M - y^i) = y^{i+1}$, then $\sigma_L(L, y^{i+1}) = y^{i+1}$ and if $\sigma_R(R, y^i) = 2M - y^{i+1}$, then $\sigma_R(R, 2M - y^{i+1}) = 2M - y^{i+1}$. Hence under consistent proposal

¹⁹It is the norm in legislative bargaining models to describe the state as solely the status quo, before a new proposer is drawn. I model the state as being described after a proposer has been drawn simply to maintain some consistency in notation with the earlier sections. This also explains the use of the redundant notation $\sigma_I(I, x)$ for party I 's strategy.

strategies a convergence path describes a lottery over equilibrium policy paths; policy dynamics are staggered and the status quo may remain unchanged for some time while the same legislator is recognised several periods in a row. When a new legislator is recognized, the status quo resumes its convergence.

Proposition 6 shows that the nonconvergence result of Proposition 4 is due to the median legislator never being recognised.

Proposition 6. *In any equilibrium in consistent proposal strategies of the legislative bargaining model, any limit point of some convergence path from state (I, x) with $x \neq M$ is bounded away from M .*

The proof shows that the main features of the results of Section 4.4, in particular those concerning convergence path payoffs under consistent strategies, can be reproduced in the legislative bargaining setting. I do not discuss the conditions for the existence of convergence paths, but these would hinge on assumptions about parties' preferences over the staggered versions of alternating outcomes. Discounting ensures that parties have a preferred such staggered alternation that is bounded away from the median and that as in my model, convergence beyond these preferred staggered alternations requires convergence paths satisfying conditions like those of (1). It is also clear that as convergence approaches the median, the effect of discounting opponents' future moderate moves outweighs the loss from parties' own present moderate moves, leading to the breakdown of convergence.

Baron (1996) characterises an equilibrium in which the median voter is indifferent between supporting the status quo and the new proposal in all periods. On the convergence paths of consistent equilibria, the median voter strictly prefers to vote against the status quo. Baron's (1996) equilibrium is in fact closely related to the equilibrium of Section 4.3. In that equilibrium, when play has reached a symmetric alternation, the median voter is indifferent between both parties' policies. Given continuation play, it would vote for any more moderate policy, since this leads to more moderate alternations, and vote against all more extreme policies.

An iid recognition rule makes it easier to verify that myopic voting is optimal for the median legislator. Consider, for example, the problematic states for myopic voting under incumbent policy persistence. Take y^i such that $\sigma_L(L, 2M - y^i) = y^{i+1}$, and consider state (R, ℓ) for $\ell \in (y^i, y^{i+1})$. R is expected to propose $r = 2M - \ell$. Suppose it deviates to $r' \in (r, 2M - y^i]$. If the median legislator supports R , policy r' is passed and in the next period the median legislator faces a lottery between a freezing of convergence at r' and a resumption of convergence by L proposing y^{i+1} . If instead it supports the status quo, in the next period the median voter faces a lottery between a freezing of convergence at ℓ and a resumption of convergence by L proposing y^{i+1} . The median legislator supports the status quo since $|M - \ell| < |M - r'|$. Since the median legislator does not affect the lottery over future proposers by its vote, it faces no cost to punish deviations.

6 Conclusion

This paper has studied the policy dynamics of a game of electoral competition between two policy-motivated parties. Although incumbent policy persistence allows opposition parties to win elections with extreme policies, an incentive to commit to more moderate policies is generated by the benefits of *imposing* moderation on future opponents. At some opportunity cost which consists of foregone policy gains in the current election, parties can, and in equilibrium do, commit to more moderate future electoral outcomes by championing moderate policies. Furthermore, since the incentives to moderate vanish as policies approach the median, it was shown that although convergence *toward* the median is a dynamically robust phenomenon, convergence *to* the median is not.

The rich policy dynamics of the model are generated by incumbent policy persistence. It is not unrealistic to suggest that incumbents and challengers are subjected to different standards by voters. In an election, incumbent politicians typically have little choice but to ‘run on their record’. Their performance in office is fresh in the minds of voters, who have had years to derive from incumbents’ decisions information about their aptitudes and preferences. Compounding this, opposition candidates or parties often elaborate and expound their platforms relative to the policies enacted by incumbents. Whatever the accepted evaluation of a politician’s or party’s term in office, they can only have marginal success in drawing voters’ attention away from it. As a consequence, their ability to propose policies to voters that differ considerably from those they championed while in office is constrained. Office-holding politicians are acutely aware of this and act accordingly. In a recent example, while less than a year into his first term, Barack Obama already frames his efforts to pass a health care reform bill through its effects on a bid for reelection which is more than three years away: ‘I intend to be president for a while and once a bill passes, I own it. And if people look and say, ‘You know what? This hasn’t reduced my costs[, ...] insurance companies are still jerking me around,’ I’m the one who’s going to be held responsible.’²⁰

References

- Acemoglu, D., Egorov, G. and K. Sonin (2008), “Dynamics and Stability of Constitutions, Coalition and Clubs”, working paper.
- Alesina, A. and G. Tabellini (1990), “A Positive Theory of Fiscal Deficits and Government Debts”, *Review of Economic Studies* 57, 403-14.
- Azzimonti, M. (2009), “Barriers to investment in polarised societies”, working paper.

²⁰‘Morning Fix: Obama, Health Care and Political Timelines’, washingtonpost.com, Monday September 14, 2009.

- Bai, J. and R. Lagunoff (2009), "On the Faustian Dynamics of Policy and Political Power", working paper.
- Banks, J. and J. Duggan (2006), "A Dynamic Model of Democratic Elections in Multidimensional Policy Spaces", working paper.
- Baron, D.P., Diermeier, D. and P. Fong (2008), "A Dynamic Theory of Parliamentary Democracy", working paper.
- Battaglini, M. and S. Coate (2007), "Inefficiency in Legislative Policymaking: A Dynamic Analysis", *American Economic Review* 97, 118-49.
- Battaglini, M. and S. Coate (2008), "A Dynamic Theory of Public Spending, Taxation and Debt", *American Economic Review* 98, 201-36.
- Bernhardt, D., Campuzano, L., Squintani, F. and O. Câmara (2009), "On the Benefits of Party Competition", *Games and Economic Behavior* 66, 685-707.
- Bernhardt, D., J. Duggan and F. Squintani (2007), "Electoral Competition with Privately informed Candidates", *Games and Economic Behavior* 58, 1-29.
- Besley, T. and S. Coate (1997), "An Economic Model of Representative Democracy", *Quarterly Journal of Economics* 112, 85-114.
- Bowen, R. and Z. Zahran (2009), "On Dynamic Compromise", working paper.
- Downs, A. (1957), *An Economic Theory of Democracy*, New York, Harper.
- Duggan, J. (2000), "Repeated Elections with Asymmetric Information", *Economics and Politics* 12, 109-35.
- Duggan, J. and M. Fey (2006), "Repeated Downsian Electoral Competition", *International Journal of Game Theory* 35, 39-69.
- Duggan, J. and T. Kalandrakis (2008), "Dynamic Legislative Policy Making", working paper.
- Fong, P. (2008), "Dynamics of Government and Policy Choice", working paper.
- Kalandrakis, T. (2004), "A Three-Player Dynamic Majoritarian Bargaining Game", *Journal of Economic Theory* 116, 294-322.
- Kalandrakis, T. (2007), "Majority Rule Dynamics with Endogenous Status Quo", working paper.
- Kalandrakis, T. (2008), "A Reputational Theory of Two Party Competition", working paper.
- Key, V.O. (1955), "A Theory of Critical Elections", *Journal of Politics* 17, 3-18.
- Key, V.O. (1958), *Politics, Parties and Pressure Groups*, 4th Edition, New York, Thomas Y. Crowell.
- Kollman, K., J. Miller and S. Page (1992), "Adaptive Parties in Spatial Elections", *American Political Science Review* 86, 929-37.
- Kollman, K., J. Miller and S. Page (1998), "Political Parties and Electoral Landscapes", *British Journal of Political Science* 28, 139-58.
- Kramer, G.H. (1977), "A Dynamical Model of Political Equilibrium", *Journal of Economic Theory* 16, 310-34.

- Laver, M. (2005), “Policy and the Dynamics of Political Competition”, *American Political Science Review* 99, 263-81.
- Lipset, S.M. and S. Rokkan (1967), Eds., *Party Systems and Voter Alignments: Cross-National Perspectives*, New York, Free Press.
- de Marchi, S. (1999), “Adaptive Models and Electoral Instability”, *Journal of Theoretical Politics* 11, 393-419.
- Miller, A.H. and M.P. Wattenberg (1985), “Throwing the Rascals Out: Policy and Performance Evaluations of Presidential Candidates, 1952-1980”, *American Political Science Review* 79, 359-72.
- Nadeau, R. and M.S. Lewis-Beck (2001), “National Economic Voting in U.S. Presidential Elections”, *The Journal of Politics* 63, 159-81.
- Osborne, M.J. (1995), “Spatial Models of Political Competition Under Plurality Rule: A Survey of Some Explanations of the Number of Candidates and the Positions they Take”, *Canadian Journal of Economics* 28, 261-301.
- Osborne, M.J. and A. Slivinski, “A Model of Political Competition with Citizen-Candidates”, *Quarterly Journal of Economics* 111, 65-96.
- Persson, T. and G. Tabellini (2000), *Political Economics*, Cambridge, MIT Press.
- Sundquist, J.L. (1983), *Dynamics of the Party System: Alignment and Realignment of Political Parties in the United States*, 2nd Edition, Washington D.C., The Brookings Institute.
- Wittman, D. (1977), “Candidates with Policy Preferences: A Dynamic Model”, *Journal of Economic Theory* 14, 180-9.
- Van Weelden, R. (2009), “Candidates and Credible Promises in Repeated Elections”, working paper.

A Appendix

Proof of Proposition 1: As noted in the text, $\frac{1}{1-\delta_J}u_J(M)$ is a subgame perfect equilibrium payoff for party J following any history. Since party J can always enforce this payoff by committing to policy M following any history, this payoff is the lowest SPE payoff for J . Hence a policy path $\{y^i\}$ is a subgame perfect equilibrium policy path only if $\sum_{i=0}^{\infty} \delta^i u_J(y^i) \geq \frac{1}{1-\delta_J}u_J(M)$ for all J and all i .

The first step in the proof shows that the game’s only subgame perfect equilibrium policy path following any history is the indefinite repetition of the median policy. Strict concavity is needed to ensure that if $y \neq M$ is strictly on party J ’s side of the median, then $u_J(y) - u_J(M) < u_{-J}(M) - u_{-J}(y)$.²¹ This holds since any strictly concave functions u_L and u_R defined on $[0, 1]$ with u_L strictly decreasing and u_R strictly increasing can be normalised such that $|u'_L(M)| = |u'_R(M)|$.

²¹Any assumptions that yields this property are sufficient for the result of Proposition 1. For example, if u_L and u_R are weakly concave but strictly concave in a neighbourhood of M .

Suppose $y < M$. By strict concavity, for all $\ell \in [y, M)$ we have $|u'_L(\ell)| < |u'_L(M)| = |u'_R(M)| < |u'_R(\ell)|$, and hence $u_L(y) - u_L(M) < u_R(M) - u_R(y)$.

Consider subgame perfect equilibrium policy path $\{y^i\}$ following some history with $y^0 \neq M$, and suppose that y^0 is on J 's side of the median. Define

$$D_J^0 = 0, \\ D_{-J}^0 = \frac{u_{-J}(M) - u_{-J}(y^1)}{\delta_{-J}}.$$

For any $i \geq 1$ and y^i (weakly) on J 's side of the median, define D_J^i and D_{-J}^i recursively as

$$D_J^i = \max \left\{ 0, \frac{D_J^{i-1} + [u_J(M) - u_J(y^i)]}{\delta_J} \right\}, \\ D_{-J}^i = \frac{D_{-J}^{i-1} + [u_{-J}(M) - u_{-J}(y^i)]}{\delta_{-J}}.$$

That is, interpret $D_J^i \geq 0$ as the payoff 'debt' for party J at stage i of subgame perfect equilibrium policy path $\{y^i\}$ relative to path $(M, M, \dots)$. This debt collects all deviations from payoff $u_J(M)$; if party J makes a loss with respect to $u_J(M)$ at y^i , then the equilibrium payoff from y^{i+1} needs to yield an excess of at least D_J^i over $\frac{1}{1-\delta_J}u_J(M)$. Debts grow by factor $\frac{1}{\delta_J}$ each period since they are incurred in the current period and reimbursed in later periods. Negative debts are never incurred since party J must be guaranteed the payoff $\frac{1}{1-\delta_J}u_J(M)$ after all histories.

Since $y^0 \neq M$, debts (D_L^0, D_R^0) are such that $D_J^0 > 0$ for some J . Suppose without loss of generality that $\delta_L \leq \delta_R$. First note that for all $i > 0$, it cannot be that $D_L^i = D_R^i = 0$, since $D_J^0 > 0$ and whenever $D_J^i < D_J^{i-1}$, it must be that y^i is strictly on J 's side of the median and hence that $D_{-J}^i > D_{-J}^{i-1}$. Next, note that for all J , we have that $\liminf_{i \rightarrow \infty} D_J^i = 0$, and also that $D_J^i = 0$ infinitely often. To see this, suppose that there exists some k such that $D_J^i > 0$ for all $i \geq k$. Then the equilibrium value to party J from subgame perfect equilibrium policy path $\{y^i\}_{i=k}^\infty$ is strictly less than $\frac{1}{1-\delta_J}u_J(M)$, a contradiction.

Suppose now that $y^0 < M$, and hence that $D_L^0 = 0 < D_R^0$. Then either

- i. $D_L^i = 0$ for all $i > 0$.
- ii. $D_L^i > 0$ for some $i > 0$.

In case i, it must be that $y^i \leq M$ for all $i > 0$, and hence that $\lim_{i \rightarrow \infty} D_R^i \geq \lim_{i \rightarrow \infty} \frac{D_R^0}{\delta_R^i} = \infty$, a contradiction. We now see that assuming $y^0 < M$ is without loss of generality. First, any subgame perfect equilibrium policy path that deviates from the median policy after some history must have some subsequence that begins at stage k with debt levels $D_J^k = 0 < D_{-J}^k$. Second, assume instead that $D_L^0 > 0 = D_R^0$. Then either $D_R^i = 0$ for all i , which leads to contradiction, or there exists k such that $D_L^k = 0$, in which case we must have $D_R^k > 0$. Now consider case ii

above. There must exist $n > m \geq 0$ with $n - m > 1$ such that $D_R^m > 0$, $D_L^m = D_L^n = 0$ and $D_L^i > 0$ for $i \in \{m+1, \dots, n-1\}$. We want to show that $D_R^m < D_R^n$. Consider the sequence $\{\hat{y}^i\}_{i=m+1}^n$ that solves the following minimisation problem.

$$\min_{\{y^i\}_{i=m+1}^n \in X^{n-m}} D_R^n \quad \text{subject to } D_L^m = D_L^n = 0, \text{ given } D_R^m > 0. \quad (2)$$

$\{\hat{y}^i\}_{i=m+1}^n$ exists since D_R^n is continuous and X^{n-m} is compact. Suppose that $\{\hat{y}^i\}_{i=m+1}^n$ is such that $\hat{D}_L^{n-1} > 0$, where $\hat{D}_J^i$ is the debt of party J under $\{\hat{y}^i\}_{i=m+1}^n$. Hence since $D_L^n = 0$ it must be that $\hat{y}^n < M$. Suppose that $\hat{D}_R^{n-2} + [u_R(M) - u_R(\hat{y}^{n-1})] < 0$, which implies that $\hat{D}_R^{n-1} = 0$ and that $\hat{y}^{n-1} > M$. For $\epsilon > 0$, consider $\bar{y}^{n-1} = \hat{y}^{n-1} - \epsilon$ and $\bar{y}^n = \hat{y}^n + \eta_\epsilon$, where η_ϵ is chosen such that $\bar{D}_L^n = 0$. For sufficiently small ϵ , we have that $\bar{D}_R^{n-1} = \hat{D}_R^{n-1} = 0$ and $\bar{D}_R^n < \hat{D}_R^n$, a contradiction. Now suppose that $\hat{D}_R^{n-2} + [u_R(M) - u_R(\hat{y}^{n-1})] \geq 0$. $\hat{D}_R^n$ is strictly increasing in $\hat{y}^{n-1}$ if

$$-\frac{u'_R(\hat{y}^{n-1})}{\delta_R^2} - \frac{u'_R(\hat{y}^n)}{\delta_R} \frac{d\hat{y}^n}{d\hat{y}^{n-1}} > 0, \quad (3)$$

where $\frac{d\hat{y}^n}{d\hat{y}^{n-1}}$ is given by

$$\frac{u'_L(\hat{y}^{n-1})}{\delta_L^2} - \frac{u'_L(\hat{y}^n)}{\delta_L} \frac{d\hat{y}^n}{d\hat{y}^{n-1}} = 0,$$

or $\frac{d\hat{y}^n}{d\hat{y}^{n-1}} = -\frac{1}{\delta_L} \frac{u'_L(\hat{y}^{n-1})}{u'_L(\hat{y}^n)}$, which comes from partially differentiating the constraint $D_L^n = 0$ with respect to $\hat{y}^{n-1}$ and $\hat{y}^n$. We can rewrite (3) as

$$\frac{u'_L(\hat{y}^{n-1})}{u'_L(\hat{y}^n)} > \frac{\delta_L}{\delta_R} \frac{u'_R(\hat{y}^{n-1})}{u'_R(\hat{y}^n)}.$$

Say $\hat{y}^{n-1} \geq M$. Then $|u'_L(\hat{y}^{n-1})| \geq |u'_R(\hat{y}^{n-1})|$, $\frac{\delta_L}{\delta_R} \leq 1$ and $|u'_L(\hat{y}^n)| < |u'_R(\hat{y}^n)|$ (since $\hat{y}^n < M$) imply that (3) holds, and hence that $\{\hat{y}^i\}_{i=m+1}^n$ does not solve (2), a contradiction. Hence it must be that $\hat{y}^{n-1} < M$.

This pairwise necessary condition for optimality can be used all along the sequence $\{\hat{y}^i\}_{i=m+1}^n$ to show that a solution to (2) with $\hat{y}^n < M$ must have $\hat{y}^i < M$ for all $i \in \{m+1, \dots, n-1\}$. But consider instead sequence $\{\tilde{y}^i\}_{i=m+1}^n$ with $\tilde{y}^i = M$ for all i . This sequence satisfies the constraints of (2), and is such that $\tilde{D}_R^n = \frac{D_R^m}{\delta_R^{n-m}} < D_R^n$ for any $\{y^i\}_{i=m+1}^n$ with $D^{n-1} < M$. Hence, for the purported equilibrium sequence from above, we have as desired that $D_R^n > D_R^m$. Considering the full policy sequence, we have that whenever $D_L^i > 0$ for $i \in \{m+1, n-1\}$, then $D_R^n > D_R^m$. Furthermore, whenever $D_L^i = 0$ for $i \in \{m+1, n-1\}$, then again $D_R^n > D_R^m$ since $D_L^i = 0$ only if $y^i \leq M$, and as shown above if $D_L^m = 0$, then $D_R^m > 0$. Hence, given the SPE path $\{y^i\}$ following some history for which $D_R^k > 0$, we have that $\lim_{i \rightarrow \infty} D_R^i = \infty$, a contradiction.

The previous argument shows that the unique SPE policy path following any history is $(M, M, \dots)$. It remains to be shown that both parties' strategies must call for them to commit to the median following any history. If party J 's strategy calls for some policy $y \neq M$ after some history, then party $-J$ must win the election with policy M . Since $y \neq M$, party $-J$ can win the election with a policy it prefers to M , say y' . Since following any deviation, party $-J$ payoffs revert to $\frac{1}{1-\delta_{-J}}(M)$, deviating to y' is profitable for $-J$. $\square$

A.1 Policy Dynamics

Proof of Proposition 2: Consider state (R, r) and policy path $\{(\ell^i, r^i)\}$ induced by (σ_L, σ_R) from (R, r) . First note that the policy path following state (R, M) can only be $(M, M, \dots)$. The following lemma, which implies the rest of point i, also establishes part of point ii. It shows that if (R, r) is such that $r < M$, party L will either not participate in the election or commit to a losing policy, while if $r \geq M$ party L will defeat party R .

Lemma 2. *In any MPE (σ_L, σ_R) , $\sigma_L(R, r) \in X \setminus W(R, r) \cup \{Out\}$ for all $r < M$, and $\sigma_R(L, \ell) \in X \setminus W(L, \ell) \cup \{Out\}$ for all $\ell > M$. Furthermore, $\sigma_L(R, r) \leq M$ for $r > M$ and $\sigma_R(L, \ell) \geq M$ for $\ell < M$.*

Proof: Consider some MPE (σ_L, σ_R) with $\sigma_L(R, r) \in [r, 2M - r]$ for some $r < M$. Consider a one-shot deviation by L at state (R, r) to Out . The payoff to this deviation is

$$u_L(r) + \delta_L V_L(\sigma_L, \sigma_R; (R, r)),$$

while the payoff to $\sigma_L(R, r)$ is $V_L(\sigma_L, \sigma_R; (R, r))$. Hence the deviation is unprofitable if and only if

$$V_L(\sigma_L, \sigma_R; (R, r)) \geq \frac{1}{1 - \delta_L} u_L(r). \quad (4)$$

Since $r < M$, the policy path following (R, r) most favourable to L is $(r, r, \dots)$. Hence we have that

$$V_L(\sigma_L, \sigma_R; (R, r)) \leq \frac{1}{1 - \delta_L} u_L(r). \quad (5)$$

(4) and (5) imply that $V_L(\sigma_L, \sigma_R; (R, r)) = \frac{1}{1 - \delta_L} u_L(r)$, which holds if and only if $\sigma_L(R, r) = r$ and $\sigma_R(L, r) = r$. Now consider a deviation for R in state (ℓ, r) to $r^d \in (r, 2M - r]$. Any policy path $\{(\ell^i, r^i)\}$ induced by (σ_L, σ_R) from (R, r^d) must be such that $\ell^i > r$ and $r^i > r$ for all i .

Hence the payoff to r^d is

$$\begin{aligned}
u_R(r^d) + \sum_{i=0}^{\infty} \delta_R^{2i+1} [u_R(\ell^i) + \delta_R u_R(r^i)] \\
&> u_R(r^d) + \frac{\delta_R}{1 - \delta_R^2} [u_R(\ell) + u_R(r)] \\
&> \frac{1}{1 - \delta_R^2} u_R(r).
\end{aligned}$$

a contradiction.

For the second claim, take (R, r) with $r > M$ and $\sigma_L(R, r) > M$. Consider a deviation to some $\ell^d \in (M, \sigma_L(R, r))$. By the first part of the lemma, the payoff to ℓ^d is given by

$$\begin{aligned}
\frac{1}{1 - \delta_L} u_L(\ell^d) &> \frac{1}{1 - \delta_L} u_L(\sigma_L(R, r)) \\
&= V_L(\sigma_L, \sigma_R; (R, r)),
\end{aligned}$$

a contradiction. Now take (R, r) with $r > M$ and $\sigma_L(R, r) = \text{Out}$. Consider deviation to some $\ell^d \in (M, r)$. By the first part of the lemma, the payoff to ℓ^d is given by

$$\begin{aligned}
\frac{1}{1 - \delta_L} u_L(\ell^d) &> \frac{1}{1 - \delta_L} u_L(r) \\
&= V_L(\sigma_L, \sigma_R; (R, r)).
\end{aligned}$$

a contradiction. □

For point ii of Proposition 2, note that by Lemma 2 the sequence $\{\ell^i\}$ is a weakly increasing sequence of real numbers bounded by ℓ^0 and M , and hence converges to some limit $\hat{\ell}$. The sequence $\{r^i\}$ is a weakly decreasing sequence of real numbers bounded by M and r^0 , and hence $\{r^i\} \rightarrow \hat{r}$. Furthermore, it must be that $\hat{\ell} = 2M - \hat{r}$. Suppose instead that $\hat{\ell} - (2M - \hat{r}) = \epsilon > 0$. Consider $n \in \mathbf{N}$ such that $\hat{\ell} - \ell^i < \epsilon$ for all $i \geq n$. Then for $j \geq N$

$$\begin{aligned}
2M - \ell^j &< 2M - \hat{\ell} + \epsilon \\
&= \hat{r} \\
&\leq r^{j+1}
\end{aligned}$$

and hence $r^{j+1} \notin W(L, \ell^j)$ and there can be no $\sigma_R(L, \ell^j)$ such that $\tau((L, \ell^j), \sigma_R(L, \ell^j)) = r^{j+1}$, a contradiction. A similar argument shows that it cannot be that $\hat{\ell} < 2M - \hat{r}$. Hence $\hat{r} = 2M - \hat{\ell}$.

To complete the proof of Proposition 2, it remains to be shown that $\sigma_L(R, \hat{r}) = \hat{\ell}$ and $\sigma_R(L, \hat{\ell}) = \hat{r}$. Suppose first that $(\ell^i, r^i) = (\hat{\ell}, \hat{r})$ for some i . Then $(\ell^j, r^j) = (\hat{\ell}, \hat{r})$ for all $j > i$ and

it must be that $\sigma_L(R, \hat{r}) = \hat{\ell}$ and $\sigma_R(L, \hat{\ell}) = \hat{r}$. Suppose now that $(\ell^i, r^i) \neq (\hat{\ell}, \hat{r})$ for all i , and that $\sigma_R(L, \hat{\ell}) = r < \hat{r}$. Consider $\Delta > 0$ such that

$$u_L(\hat{\ell}) + \frac{\delta_L}{1 - \delta_L^2} \Psi_L^-(2M - r) > \frac{1}{1 - \delta_L^2} \Psi_L^+(\hat{\ell}) + \Delta. \quad (6)$$

Such a Δ exists by Lemma 1 since $r < \hat{r}$. Since u_L is continuous and $\{\ell^i\} \rightarrow \hat{\ell}$, there exists $n \in \mathbf{N}$ and $\epsilon > 0$ such that for all $i \geq n$, $\hat{\ell} - \ell^i < \epsilon$ and $u_L(\ell^i) - u_L(\hat{\ell}) < \Delta$. Now, for any $j \geq N$,

$$\begin{aligned} V_L(\sigma_L, \sigma_R; (R, r^j)) &= u_L(\ell^j) + \sum_{i=0}^{\infty} \delta_L^{2i+1} [u_L(r^{j+1+i}) + \delta_L u_L(\ell^{j+1+i})] \\ &\leq u_L(\ell^j) + \sum_{i=0}^{\infty} \delta_L^{2i+1} \Psi_L^-(r^{j+1+i}) \\ &\leq u_L(\ell^j) + \frac{\delta_L}{1 - \delta_L^2} \Psi_L^-(\hat{\ell}) \\ &< \frac{1}{1 - \delta_L^2} \Psi_L^+(\hat{\ell}) + \Delta. \end{aligned} \quad (7)$$

The first inequality follows from the fact that $\ell^{j+1+i} \geq 2M - r^{j+1+i}$ for all i . The second inequality follows by Lemma 1 from the fact that $r^{j+1+i} \geq \hat{r}$ for all i . In state (R, r^j) , consider a deviating strategy by L , σ_L^d , with the properties

$$\begin{aligned} \sigma_L^d(R, r^j) &= \hat{\ell} \text{ and} \\ \sigma_L^d(R, r') &= 2M - r' \text{ for all } r' \leq r. \end{aligned}$$

Consider the policy path $\{(\ell^i, r^i)\}$ induced by (σ_L^d, σ_R) from (R, r^j) . Note that $r^0 = r$ and $\ell^i = 2M - r^i$. The payoff to σ_L^d is

$$\begin{aligned} u_L(\hat{\ell}) + \sum_{i=0}^{\infty} \delta_L^{2i+1} \Psi_L^-(2M - r^i) &\geq u_L(\hat{\ell}) + \frac{\delta_L}{1 - \delta_L^2} \Psi_L^-(2M - r) \\ &> \frac{1}{1 - \delta_L^2} \Psi_L^+(\hat{\ell}) + \Delta \\ &> V_L(\sigma_L, \sigma_R; (R, r^j)), \end{aligned}$$

a contradiction. The first inequality follows from Lemma 1 and the fact that $r^i \leq r$ for all i , the second from (6) and the third from (7). The same proof applies to show that $\sigma_L(R, \hat{r}) = \hat{\ell}$. This completes the proof of Proposition 2. $\square$

A.2 Bounded Extremism

The following lemma provides a lower bound on equilibrium payoffs.

Lemma 3. *Consider MPE (σ_L, σ_R) . In state (R, r) with $r > M$, the payoff to party L from policy $\ell \in W(R, r)$ for some $\ell \leq M$ is at least $\frac{1}{1-\delta_L^2} \Psi_L^+(\ell)$. Furthermore,*

$$V_L(\sigma_L, \sigma_R; (R, r)) \geq \frac{1}{1-\delta_L^2} \Psi_L^+(2M-r).$$

The statement for party R is symmetric.

Proof: Given state (R, r) with $r > M$, consider the strategy σ'_L for L with the properties

$$\begin{aligned} \sigma'_L(R, r) &= 2M - r \text{ and} \\ \sigma'_L(R, r') &= 2M - r' \text{ for all } r' \leq 2M - \ell^*. \end{aligned}$$

Consider the policy path $\{(\ell^i, r^i)\}$ induced by (σ'_L, σ_R) from (R, r) . The payoff to σ'_L is

$$\begin{aligned} u_L(2M-r) + \sum_{i=0}^{\infty} \delta_L^{2i+1} \Psi_L^-(2M-r^i) &\geq u_L(2M-r) + \frac{\delta_L}{1-\delta_L^2} \Psi_L^-(2M-r) \\ &= \frac{1}{1-\delta_L^2} \Psi_L^+(2M-r). \end{aligned}$$

The inequality follows by Lemma 1 since $r^i \leq r$ for all i . The rest of the lemma follows by since $2M-r \in W(R, r)$ implies that $V_L(\sigma_L, \sigma_R; (R, r)) \geq \frac{1}{1-\delta_L^2} \Psi_L^+(2M-r)$. □

The following lemma establishes the bound on the extremism of alternating long-run policy outcomes.

Lemma 4. *If policies $(\ell, 2M-\ell)$ are alternating long-run policy outcomes, then $\ell \geq \max\{\ell^*, 2M-r^*\}$.*

Proof: Suppose that $\ell^* \geq 2M-r^*$. Suppose (ℓ, r) is a long-run policy outcome under MPE (σ_L, σ_R) and $\ell < \ell^*$. By Lemma 3, party L can obtain a payoff of at least $\frac{1}{1-\delta_L^2} \Psi_L^+(\ell^*)$ by committing to ℓ^* in state (R, r) . However, $V_L(\sigma_L, \sigma_R; (R, r)) = \frac{1}{1-\delta_L^2} \Psi_L^+(\ell) < \frac{1}{1-\delta_L^2} \Psi_L^+(\ell^*)$ by Lemma 1 since $\ell < \ell^*$, a contradiction. □

The following lemma, an immediate consequence of Markov strategies, is stated here without proof for the sake of clarity.

Lemma 5. (Value to One-Shot Deviations) *Consider strategy pair (σ_L, σ_R) . Consider one-shot deviation to ℓ^d for L at (R, r) such that $\ell^d = \sigma_L(R, r')$ for $r' \neq r$. Then the payoff to party L from ℓ^d is $V_L(\sigma_L, \sigma_R; (R, r'))$.*

The following lemma verifies the equilibrium construction of Section 4.3.

Lemma 6. *If $\ell^* \geq 2M - r^*$, the strategy profile $(\sigma_L^{\ell^*}, \sigma_R^{my})$ forms an equilibrium. If $\ell^* < 2M - r^*$, the strategy profile $(\sigma_L^{my}, \sigma_R^*)$ forms an equilibrium.*

Proof: Suppose that $\ell^* \geq 2M - r^*$. First verify the optimality of L 's proposed strategy. Given $(\sigma_L^{\ell^*}, \sigma_R^{my})$ compute

$$V_L(\sigma_L^{\ell^*}, \sigma_R^{my}; (R, r)) = \begin{cases} \frac{1}{1-\delta_L^2} \Psi_L^+(\ell^*) & \text{for } r \in [2M - \ell^*, 1], \\ \frac{1}{1-\delta_L^2} \Psi_L^+(2M - r) & \text{for } r \in [M, 2M - \ell^*), \\ \frac{1}{1-\delta_L^2} u_L(r) & \text{for } r \in [0, M). \end{cases}$$

Note that for all r, r' such that $r > r'$, $\sigma_L(R, r) \in W(R, r)$ and $\sigma_L(R, r) \neq \sigma_L(R, r') \in W(R, r')$,

$$V_L(\sigma_L^{\ell^*}, \sigma_R^{my}; (R, r)) > V_L(\sigma_L^{\ell^*}, \sigma_R^{my}; (R, r')).$$

Hence, by Lemma 5, at any state (R, r) such that $\sigma_L(R, r) \in W(R, r)$, party L cannot profit from one-shot deviation ℓ^d such that $\sigma_L(R, r') = \ell$ for some $r' \neq r$. Hence only one-shot deviations $\ell^d \in [0, \ell^*) \cup (M, 1]$ can be profitable for L at some state.

The value of setting $\ell^d \in [0, \ell^*)$ if winning at (R, r) is

$$\begin{aligned} & u_L(\ell^d) + \delta_L u_L(2M - \ell^d) + \delta_L^2 V_L(\sigma_L^{\ell^*}, \sigma_R^{my}; (R, 2M - \ell^d)) \\ &= \Psi_L^+(\ell^d) + \frac{\delta_L^2}{1 - \delta_L^2} \Psi_L^+(\ell^*). \end{aligned}$$

$\ell^d \in [0, \ell^*)$ is winning only in states (R, r) with $r \in [2M - \ell^d, 1] \cup [0, \ell^d]$. For $r \in [2M - \ell^d, 1]$

$$\begin{aligned} V_L(\sigma_L^{\ell^*}, \sigma_R^{my}; (R, r)) &= \frac{1}{1 - \delta_L^2} \Psi_L^+(\ell^*) \\ &> \Psi_L^+(\ell^d) + \frac{\delta_L^2}{1 - \delta_L^2} \Psi_L^+(\ell^*), \end{aligned}$$

where the inequality follows from Lemma 1 since $\ell^d < \ell^*$. For $r \in [0, \ell^d]$

$$\begin{aligned} V_L(\sigma_L^{\ell^*}, \sigma_R^{my}; (R, r)) &= \frac{1}{1 - \delta_L^2} u_L(r) \\ &> \Psi_L^+(\ell^d) + \frac{\delta_L^2}{1 - \delta_L^2} \Psi_L^+(\ell^*), \end{aligned}$$

where the inequality follows since $r \leq \ell^d$.

The value of setting $\ell^d \in (M, 1]$ if winning at (R, r) is

$$\frac{1}{1 - \delta_L^2} u_L(\ell^d).$$

$\ell^d \in (M, 1]$ is winning only in states (R, r) with $r \in [2M - \ell^d, M] \cup [\ell^d, 1]$. For $r \in [2M - \ell^d, M]$

$$\begin{aligned} V_L(\sigma_L^{\ell^*}, \sigma_R^{my}; (R, r)) &= \frac{1}{1 - \delta_L^2} u_L(r) \\ &> \frac{1}{1 - \delta_L^2} u_L(\ell^d), \end{aligned}$$

where the inequality follows since $r < \ell^d$. For $r \in [\ell^d, 1]$

$$\begin{aligned} V_L(\sigma_L^{\ell^*}, \sigma_R^{my}; (R, r)) &> \frac{1}{1 - \delta_L^2} u_L(M) \\ &> \frac{1}{1 - \delta_L^2} u_L(\ell^d), \end{aligned}$$

where the first inequality follows since $r > M$, and the second since $\ell^d > M$. Hence, no profitable deviation for L exists and $\sigma_L^{\ell^*}$ is optimal when facing σ_R^{my} .

Now verify the optimality of R 's proposed strategy. Given $(\sigma_L^{\ell^*}, \sigma_R^{my})$ compute

$$V_R(\sigma_R^{my}, \sigma_L^{\ell^*}; (L, \ell)) = \begin{cases} u_R(2M - \ell) + \frac{\delta_R}{1 - \delta_R^2} \Psi_R^-(\ell^*) & \text{for } \ell \in [0, \ell^*), \\ \frac{1}{1 - \delta_R^2} \Psi_R^+(\ell) & \text{for } \ell \in [\ell^*, M), \\ \frac{1}{1 - \delta_R^2} u_R(\ell) & \text{for } \ell \in [M, 1). \end{cases}$$

Again, note that for all ℓ, ℓ' such that $\ell < \ell'$, $\sigma_R(L, \ell) \in W(L, \ell)$ and $\sigma_R(L, \ell) \neq \sigma_R(L, \ell') \in W(L, \ell')$

$$V_R(\sigma_R^{my}, \sigma_L^{\ell^*}; (L, \ell)) > V_R(\sigma_R^{my}, \sigma_L^{\ell^*}; (L, \ell')).$$

Hence, by Lemma 5, at any state (L, ℓ) such that $\sigma_R(L, \ell) \in W(L, \ell)$, party R cannot profit by deviating to any r^d such that $\sigma_R(L, \ell') = r^d$ for some $\ell' \neq \ell$. Hence only one-shot deviations $r^d \in [0, M)$ can be profitable for R at some state. That these cannot be profitable for R follows from a verification similar to that for deviations $\ell^d \in (M, 1]$ for L above. Hence, no profitable deviation for R exists and σ_R^{my*} is optimal when facing $\sigma_L^{\ell^*}$. □

A.3 Consistent Markov Perfect Equilibria

The following Lemma characterises convergence paths under consistent strategies.

Lemma 7. *Consider consistent Markov strategies σ_L and σ_R .*

- i. If $\sigma_L(R, r) = \ell \in (\max\{2M - r, 0\}, M]$ for some $r > M$, then $\sigma_L(R, r') = \ell$ for all $r' \in [2M - \ell, r)$.*
- ii. Suppose (σ_L, σ_R) form a consistent equilibrium. If $\sigma_L(R, r) = \ell \in (\max\{2M - r, 0\}, M]$ for some $r > M$, then $\sigma_R(L, \ell') = 2M - \ell'$ for all $\ell' \in [\max\{2M - r, 0\}, \ell)$.*

Both statements for R are symmetric.

Proof: Part i is immediate from the definition of consistent Markov strategies. For part ii, consider consistent equilibrium (σ_L, σ_R) , $r > M$ and $\sigma_L(R, r) = \ell > \max\{2M - r, 0\}$. Suppose for some $\ell' \in [\max\{2M - r, 0\}, \ell)$, $\sigma_R(L, \ell') = r' < 2M - \ell'$. There are two cases. First, suppose that $r' \geq 2M - \ell$. Consider the one-shot deviation by R to $2M - \ell'$ in state (L, ℓ') . The payoff to this deviation is

$$\begin{aligned} & u_R(2M - \ell') + \delta_R V_R(\sigma_R, \sigma_L; (R, 2M - \ell')) \\ &= u_R(2M - \ell') + \delta_R V_R(\sigma_R, \sigma_L; (R, r')) \\ &> u_R(r') + \delta_R V_R(\sigma_R, \sigma_L; (R, r')) \\ &= V_R(\sigma_R, \sigma_L; (L, \ell')). \end{aligned}$$

a contradiction. The second line follows since $\sigma_L(R, r') = \ell$ for all $r' \in [2M - \ell, r]$, and the inequality since $r' < 2M - \ell'$.

Second, suppose $r' < 2M - \ell$. Then by the part i of the lemma it must be that $\sigma_R(L, \ell'') = r'$ for all $\ell'' \in [\ell', 2M - r']$. By reversing the roles in the proof of the first case above, it can be seen that L can profitably deviate to $2M - r'$ at (R, r') . □

The following lemma characterises payoffs on consistent equilibrium convergence paths.

Lemma 8. *Consider long-run policy outcome $(\hat{\ell}, 2M - \hat{\ell})$ for some $\hat{\ell} > \max\{\ell^*, 2M - r^*\}$, associated consistent equilibrium (σ_L, σ_R) and convergence path $\{y^i\} \rightarrow \hat{\ell}$ starting from some state. Take state $(R, 2M - y^i)$ such that $\sigma_L(R, 2M - y^i) = y^{i+1}$ with $i > 0$. Then*

$$V_L(\sigma_L, \sigma_R; (R, 2M - y^i)) = \frac{1}{1 - \delta_L^2} \Psi_L^+(y^i). \quad (8)$$

Furthermore,

$$\frac{1}{1 - \delta_L^2} \Psi_L^+(y^i) = u_L(y^{i+1}) + \frac{\delta_L}{1 - \delta_L^2} \Psi_L^-(y^{i+2}). \quad (9)$$

The case of state (L, y^i) such that $\sigma_R(L, y^i) = 2M - y^{i+1}$ with $i > 0$ is symmetric.

Proof: Consider state $(R, 2M - y^i)$ such that $\sigma_L(R, 2M - y^i) = y^{i+1}$ with $i > 0$. Since $\hat{\ell} > \max\{\ell^*, 2M - r^*\}$, we have that $y^i < y^{i+1}$ for all i . Since $i > 0$, by Lemma 7 there exists $\epsilon > 0$ such that for all $\ell \in (y^i - \epsilon, y^i]$, $\sigma_R(L, \ell) = 2M - y^i$. For any $\bar{\epsilon} \in (0, \epsilon)$, consider one-shot deviation by L at $(R, 2M - y^i + \bar{\epsilon})$ to $y^{i+1} = \sigma_L(R, 2M - y^i)$. By Lemma 5, the value to this deviation is given by

$$\begin{aligned} V_L(\sigma_L, \sigma_R; (R, 2M - y^i)) &\leq V_L(\sigma_L, \sigma_R; (R, 2M - y^i + \bar{\epsilon})) \\ &= u_L(y^i - \epsilon) + \delta_L u_L(2M - y^i) \\ &\quad + \delta_L^2 V_L(\sigma_L, \sigma_R; (R, 2M - y^i)), \end{aligned}$$

where the inequality follows from equilibrium. This yields

$$V_L(\sigma_L, \sigma_R; (R, 2M - y^i)) \leq \frac{1}{1 - \delta_L^2} [u_L(y^i - \bar{\epsilon}) + \delta_L u_L(2M - y^i)]$$

for any $\bar{\epsilon} \in (0, \epsilon)$, and hence

$$V_L(\sigma_L, \sigma_R; (R, 2M - y^i)) \leq \frac{1}{1 - \delta_L^2} \Psi_L^+(y^i).$$

Lemma 3 yields the opposite inequality and hence

$$V_L(\sigma_L, \sigma_R; (R, 2M - y^i)) = \frac{1}{1 - \delta_L^2} \Psi_L^+(y^i).$$

The final claim of the lemma follow since

$$\begin{aligned} V_L(\sigma_L, \sigma_R; (R, 2M - y^i)) &= u_L(y^{i+1}) + \delta_L u_L(2M - y^{i+2}) \\ &\quad + \delta_L^2 V_L(\sigma_L, \sigma_R; (R, 2M - y^{i+2})). \end{aligned} \tag{10}$$

□

A.4 Bounded Moderation

A.4.1 Construction of Bound on Long-run Moderation

Lemma 9. Define mappings $\alpha_L : [\max\{\ell^*, 2M - r^*\}, M] \rightarrow (0, 1]$ and $\alpha_R : [\max\{\ell^*, 2M - r^*\}, M] \rightarrow (0, 1]$ such that

$$\begin{aligned} \frac{u'_L(\ell)}{u'_L(2M - \ell)} &= \frac{\delta_L}{\delta_L^2 + \alpha_L(\ell)(1 - \delta_L^2)} \text{ and} \\ \frac{u'_R(2M - \ell)}{u'_R(\ell)} &= \frac{\delta_R}{\delta_R^2 + \alpha_R(\ell)(1 - \delta_R^2)}. \end{aligned} \tag{11}$$

Define ℓ^{**} such that $\alpha_L(\ell^{**}) + \alpha_R(\ell^{**}) = 1$. Then α_L , α_R and ℓ^{**} are well-defined and furthermore $\ell^{**} \in (\max\{\ell^*, 2M - r^*\}, M)$.

Proof: First note that since u_L is concave $\frac{u'_L(\ell)}{u'_L(2M-\ell)}$ is strictly increasing in $\ell \in [\ell^*, M]$, with a minimum of δ_L and a maximum of 1. Now $\frac{\delta_L}{\delta_L^2 + \alpha_L(1-\delta_L^2)}$ is strictly decreasing in $\alpha_L \in [0, 1]$, with a minimum of δ_L and a maximum of $\frac{1}{\delta_L}$. $\alpha_L(\ell)$ is well-defined for $\ell \in [\max\{\ell^*, 2M - r^*\}, M]$ since $\frac{u'_L(\max\{\ell^*, 2M - r^*\})}{u'_L(2M - \max\{\ell^*, 2M - r^*\})} \geq \delta_L$. Also, $\alpha_L(\ell) \in (0, 1]$ for all ℓ since $\alpha_L(M) = \frac{\delta_L}{1+\delta_L}$ and $\alpha_L(\ell^*) = 1$. Similarly, $\alpha_R(\ell)$ is well-defined. Furthermore, $\alpha_L(\ell) + \alpha_R(\ell)$ is strictly decreasing in $\ell \in [\max\{\ell^*, 2M - r^*\}, M]$, with $\alpha_L(M) + \alpha_R(M) < 1$ and $\alpha_L(\max\{\ell^*, 2M - r^*\}) + \alpha_R(\max\{\ell^*, 2M - r^*\}) > 1$. Thus $\ell^{**} \in (\max\{\ell^*, 2M - r^*\}, M)$. $\square$

To understand the derivation of α_L and α_R , consider $y^i, y^{i+2} = y^i + \Delta$ for some $\Delta > 0$ and $\alpha_L \in [0, 1]$ such that

$$\frac{1}{1 - \delta_L^2} \Psi_L^+(y^i) = u_L(y^i + \alpha_L \Delta) + \frac{\delta_L}{1 - \delta_L^2} \Psi_L^-(y^i + \Delta). \quad (12)$$

α_L is well-defined since evaluating (12) at $\alpha_L = 0$ yields

$$\begin{aligned} \frac{1}{1 - \delta_L^2} \Psi_L^+(y^i) &= u_L(y^i) + \frac{\delta_L}{1 - \delta_L^2} \Psi_L^-(y^i) \\ &< u_L(y^i) + \frac{\delta_L}{1 - \delta_L^2} \Psi_L^-(y^i + \Delta), \end{aligned}$$

while evaluating (12) at $\alpha_L = 1$ yields

$$\begin{aligned} \frac{1}{1 - \delta_L^2} \Psi_L^+(y^i) &> \frac{1}{1 - \delta_L^2} \Psi_L^+(y^i + \Delta) \\ &= u_L(y^i + \Delta) + \frac{\delta_L}{1 - \delta_L^2} \Psi_L^-(y^i + \Delta), \end{aligned}$$

where both inequalities follow from Lemma 1. The limit of (12) as $\Delta \rightarrow 0$ yields that α_L is determined by (11) evaluated at y^i .

The following lemma establishes the bound on the moderation of robust long-run policy outcomes.

Lemma 10. *If policies $(\hat{\ell}, 2M - \hat{\ell})$ are robust long-run policy outcome under some consistent equilibrium, then $\hat{\ell} \leq \ell^{**}$.*

Proof: The following lemma establishes the properties of the recursive equation (1) that determine consistent equilibrium convergence path policies that allow us to determine possible convergence points.

Lemma 11. *Consider robust long-run policy outcomes $(\hat{\ell}, 2M - \hat{\ell})$ under consistent equilibrium (σ_L, σ_R) and associated convergence path $\{y^i\}$ starting from some state.*

i. Suppose that

$$\frac{u'_L(\hat{\ell})}{u'_L(2M - \hat{\ell})} < \frac{\delta_L}{\delta_L^2 + \alpha_L(1 - \delta_L^2)} \quad (13)$$

for some $\alpha_L \in [0, 1]$ and that $\sigma_L(R, 2M - y^{i-1}) = y^i$ for some i . Then $y^i - y^{i-1} > \frac{\alpha_L}{1 - \alpha_L}(y^{i+1} - y^i)$.

ii. Conversely, suppose that

$$\frac{u'_L(y^j)}{u'_L(2M - y^j)} > \frac{\delta_L}{\delta_L^2 + \alpha_L(1 - \delta_L^2)} \quad (14)$$

for some $\alpha_L \in [0, 1]$ and that $\sigma_L(R, 2M - y^{j-1}) = y^j$. Then $y^i - y^{i-1} < \frac{\alpha_L}{1 - \alpha_L}(y^{i+1} - y^i)$ for all $i \geq j$.

The case for party R is symmetric.

Proof: To prove part i of the lemma, first prove the following claim. Suppose that for some $\alpha_L \in [0, 1]$ and y, Δ such that $y - \Delta \in [\ell^*, M]$

$$\Psi_L^+(y - \Delta) - \Psi_L^+(y - (1 - \alpha_L)\Delta) \leq \delta_L[\Psi_L^-(y) - \Psi_L^-(y - (1 - \alpha_L)\Delta)], \quad (15)$$

then for any $y' \leq y$ and $n \in \mathbf{N}$ such that $y' - 2^n \Delta \in [\ell^*, M]$

$$\Psi_L^+(y' - 2^n \Delta) - \Psi_L^+(y' - 2^n(1 - \alpha_L)\Delta) \leq \delta_L[\Psi_L^-(y') - \Psi_L^-(y' - 2^n(1 - \alpha_L)\Delta)] \quad (16)$$

with the inequality strict if $y' \neq y$ or $n > 0$. Note that (15) implies that on an infinite convergence path for some CMPE for which $\sigma_R(L, \ell) = 2M - (y - \Delta)$, $\sigma_L(R, 2M - (y - \Delta)) - y \geq \alpha_L \Delta$. The claim states that if party R 's successive policy choices on some CMPE convergence path are $2M - (y - \Delta)$ and $2M - y$ and party L is (weakly) willing to moderate to $y - (1 - \alpha_L)\Delta$ when in state $(R, 2M - (y - \Delta))$,²² then in another CMPE convergence path in which party R 's successive policies are $2M - (y' - \Delta')$ and $2M - y'$ with $y' \leq y$, then party L is strictly willing to moderate to $y' - (1 - \alpha_L)\Delta'$ in state $(R, 2M - (y' - \Delta'))$, where $\Delta' = 2^n \Delta$ for some $n \in \mathbf{N}$.

To prove the claim, note first that, for $y' \leq y$

$$\begin{aligned} \Psi_L^+(y' - \Delta) - \Psi_L^+(y' - (1 - \alpha_L)\Delta) &\leq \Psi_L^+(y - \Delta) - \Psi_L^+(y - (1 - \alpha_L)\Delta) \\ &\leq \delta_L[\Psi_L^-(y) - \Psi_L^-(y - (1 - \alpha_L)\Delta)] \\ &\leq \delta_L[\Psi_L^-(y') - \Psi_L^-(y' - (1 - \alpha_L)\Delta)], \end{aligned}$$

with the first and third inequalities strict if $y' \neq y$. The first inequality follows from the strict concavity of Ψ_L^+ , the second from (15), and the third from the strict concavity of Ψ_L^- . Given (15), the above shows that

$$\Psi_L^+(y - 2\Delta) - \Psi_L^+(y - (2 - \alpha_L)\Delta) < \delta_L[\Psi_L^-(y - \Delta) - \Psi_L^-(y - (2 - \alpha_L)\Delta)],$$

²²That is, moderate by $\alpha_L \Delta$.

and

$$\Psi_L^+(y - (2 - \alpha_L)\Delta) - \Psi_L^+(y - 2(1 - \alpha_L)\Delta) < \delta_L[\Psi_L^-(y - (1 - \alpha_L)\Delta) - \Psi_L^-(y - 2(1 - \alpha_L)\Delta)]. \quad (17)$$

Hence we have that

$$\begin{aligned} \delta_L[\Psi_L^-(y) - \Psi_L^-(y - 2(1 - \alpha_L)\Delta)] &= \delta_L[\Psi_L^-(y) - \Psi_L^-(y - (1 - \alpha_L)\Delta)] \\ &\quad + \delta_L[\Psi_L^-(y - (1 - \alpha_L)\Delta) - \Psi_L^-(y - 2(1 - \alpha_L)\Delta)] \\ &> \Psi_L^+(y - \Delta) - \Psi_L^+(y - (1 - \alpha_L)\Delta) \\ &\quad + \Psi_L^+(y - (2 - \alpha_L)\Delta) - \Psi_L^+(y - 2(1 - \alpha_L)\Delta) \\ &> \Psi_L^+(y - 2\Delta) - \Psi_L^+(y - \Delta(2 - \alpha_L)) \\ &\quad + \Psi_L^+(y - (2 - \alpha_L)\Delta) - \Psi_L^+(y - 2(1 - \alpha_L)\Delta) \\ &= \Psi_L^+(y - 2\Delta) - \Psi_L^+(y - 2(1 - \alpha_L)\Delta). \end{aligned}$$

The first inequality follows from (15) and (17), and the second inequality follows from Lemma 1 since $y - (1 - \alpha_L)\Delta = y - \Delta(2 - \alpha_L) - (y - 2\Delta) = \alpha_L\Delta$. The claim follows by applying the above argument recursively.

To complete the proof of part i of Lemma 11, consider (13). This condition guarantees that for arbitrarily small Δ , party L is willing to take up share $\alpha_L\Delta$ of moderation Δ from $y - \Delta$ to y . Hence, there exists some $\tilde{\Delta}$ such that for all $\Delta < \tilde{\Delta}$,

$$\Psi_L^+(\hat{\ell} - \Delta) - \Psi_L^+(\hat{\ell} - (1 - \alpha_L)\Delta) < \delta_L[\Psi_L^-(\hat{\ell}) - \Psi_L^-(\hat{\ell} - (1 - \alpha_L)\Delta)].$$

Thus, by the earlier claim, for all $y < \hat{\ell}$ and Δ such that $y - \Delta > \ell^*$,

$$\Psi_L^+(y - \Delta) - \Psi_L^+(y - (1 - \alpha_L)\Delta) < \delta_L[\Psi_L^-(y) - \Psi_L^-(y - (1 - \alpha_L)\Delta)].$$

This implies that for y^i such that $\sigma_L(R, 2M - y^{i-1}) = y^i$, $y^i - y^{i-1} > \frac{\alpha_L}{1 - \alpha_L}(y^{i+1} - y^i)$.

The proof of part ii of Lemma 11 follows along the lines of part i. While part i is backward-looking, part ii is forward-looking. That is, part i establishes that if at the limit point of a consistent equilibrium convergence path party L is willing to undertake share α_L of all marginal moderations, then it was also willing to undertake share α_L of all past moderate moves. In contrast, part ii shows that if at some point on a convergence path, party L would be unwilling to undertake share α_L of marginal moderations, then it will undertake less than share α_L of all future moderations on the convergence path. Evidently, part ii is useful to establish conditions for nonconvergence, while part i helps establish conditions for convergence. $\square$

Now for the proof of Lemma 10, consider a robust long-run policy outcome $(\hat{\ell}, 2M - \hat{\ell})$ with $\hat{\ell} > \ell^{**}$ and associated CMPE (σ_L, σ_R) . Consider state (R, r) with $2M - r < \hat{\ell}$ and

convergence path $\{y^i\} \rightarrow \hat{\ell}$ given (R, r) with $\sigma_L(R, 2M - y^0) = y^1$. Fix n such that $y^n > \ell^{**}$ and $\sigma_L(R, 2M - y^n) = y^{n+1}$. Hence

$$\begin{aligned} \frac{u'_L(y^n)}{u'_L(2M - y^n)} &> \frac{u'_L(\ell^{**})}{u'_L(2M - \ell^{**})} \\ &= \frac{\delta_L}{\delta_L^2 + \alpha_L(\ell^{**})(1 - \delta_L^2)}, \end{aligned}$$

and hence by part i of Lemma 11, for all $j \geq n$,

$$y^{j+1} - y^j < \frac{\alpha_L(\ell^{**})}{1 - \alpha_L(\ell^{**})}(y^{j+2} - y^{j+1}).$$

Similarly, if $j \geq n + 1$ and $\sigma_R(L, y^j) = 2M - y^{j+1}$ then

$$y^{j+1} - y^j < \frac{\alpha_R(\ell^{**})}{1 - \alpha_R(\ell^{**})}(y^{j+2} - y^{j+1}).$$

This yields that for all $j \geq n + 1$,

$$\begin{aligned} y^{j+1} - y^j &< \frac{\alpha_L(\ell^{**})}{1 - \alpha_L(\ell^{**})} \frac{\alpha_R(\ell^{**})}{1 - \alpha_R(\ell^{**})} (y^{j+3} - y^{j+2}) \\ &< (y^{j+3} - y^{j+2}). \end{aligned}$$

Hence the convergence path $\{y^i\} \rightarrow \hat{\ell}$ contains a nonconverging subsequence, a contradiction. $\square$

A.4.2 Bound on Long-run Moderation is Tight

Given a strictly increasing sequence $\{y^i\} \rightarrow \hat{\ell}$ with $y^0 = \ell^*$ and y^i, y^{i+1} and y^{i+2} satisfying the conditions of Lemma 8 for all $i \geq 1$, consider the following strategies

$$\begin{aligned} \sigma_{L^*}^{\hat{\ell}}(R, r) &= \begin{cases} \ell^* & \text{for all } r \geq 2M - \ell^*, \\ 2M - r & \text{for all } r \in (2M - y^i, 2M - y^{i-1}) \text{ with } i > 0 \text{ odd}, \\ y^i & \text{for all } r \in [2M - y^i, 2M - y^{i-1}] \text{ with } i > 0 \text{ even}, \\ 2M - r & \text{for all } r \in [M, 2M - \hat{\ell}], \\ Out & \text{for all } r < M. \end{cases} \\ \sigma_R^{\hat{\ell}}(L, \ell) &= \begin{cases} 2M - \ell & \text{for all } \ell < \ell^*, \\ y^i & \text{for all } \ell \in [y^{i-1}, y^i] \text{ with } i > 0 \text{ odd}, \\ 2M - \ell & \text{for all } \ell \in (y^{i-1}, y^i) \text{ with } i > 0 \text{ even}, \\ 2M - \ell & \text{for all } \ell \in [\hat{\ell}, M], \\ Out & \text{for all } \ell > M. \end{cases} \end{aligned}$$

If instead $\ell^* < 2M - r^*$, then for robust long-run policy outcome $(\hat{\ell}, 2M - \hat{\ell})$ with $\hat{\ell} > 2M - r^*$, strategies $(\sigma_L^{\hat{\ell}}, \sigma_R^{\hat{\ell}})$ can be constructed in a similar manner with the roles of the parties reversed.

Lemma 12. *Suppose that $\ell^* \geq 2M - r^*$. Given $\hat{\ell} > \ell^*$ and a strictly increasing sequence $\{y^i\} \rightarrow \hat{\ell}$ with $y^0 = \ell^*$ and y^i, y^{i+1} and y^{i+2} satisfying the conditions of Lemma 8 for all $i \geq 1$, strategies $(\sigma_L^{\hat{\ell}}, \sigma_R^{\hat{\ell}})$ form a consistent equilibrium under which $(\hat{\ell}, 2M - \hat{\ell})$ are robust long-run policy outcomes.*

The equilibrium $(\sigma_L^{\hat{\ell}}, \sigma_R^{\hat{\ell}})$ in the case of $\ell^ < 2M - r^*$ can be determined similarly.*

Proof: Suppose $\ell^* \geq 2M - r^*$. First verify the optimality of L 's proposed strategy. Given $(\sigma_L^{\hat{\ell}}, \sigma_R^{\hat{\ell}})$ and the conditions of the lemma for $\{y^i\}$, compute

$$V_L(\sigma_L^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, r)) = \begin{cases} u_L(\ell^*) + \frac{\delta_L}{1-\delta_L^2} \Psi_L^-(y^1) & \text{for } r \in [2M - \ell^*, 1], \\ u_L(2M - r) + \delta_L u_L(2M - y^i) + \frac{\delta_L^2}{1-\delta_L^2} \Psi_L^+(y^{i+1}) & \text{for } r \in (2M - y^i, 2M - y^{i-1}) \text{ with } i > 0 \text{ odd,} \\ u_L(y^i) + \frac{\delta_L}{1-\delta_L^2} \Psi_L^-(y^{i+1}) & \text{for } r \in [2M - y^i, 2M - y^{i-1}] \text{ with } i > 0 \text{ even,} \\ \frac{1}{1-\delta_L^2} \Psi_L^+(2M - r) \text{ for } r \in [M, 2M - \bar{\ell}], \\ \frac{1}{1-\delta_L^2} u_L(r) \text{ for } r \in [0, M]. \end{cases}$$

As in Lemma 4, note that for all r, r' such that $r > r'$, $\sigma_L(R, r) \in W(R, r)$ and $\sigma_L(R, r) \neq \sigma_L(R, r') \in W(R, r')$,

$$V_L(\sigma_L^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, r)) > V_L(\sigma_L^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, r')).$$

Hence, by Lemma 5, at any state (R, r) such that $\sigma_L(R, r) \in W(R, r)$, party L cannot profit by deviating to any ℓ^d such that $\sigma_L(R, r') = \ell$ for some $r' \neq r$. Hence only one-shot deviations $\ell^d \in [0, \ell^*) \cup \left(\bigcup_{i>0, i \text{ even}} [y^{i-1}, y^i] \right) \cup (M, 1]$ can be profitable for L at some state. The value to setting $\ell^d \in [0, \ell^*)$ if winning at (R, r) is

$$\Psi_L^+(\ell^d) + \delta_L^2 V_L(\sigma_L^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, 2M - \ell^d)) = \Psi_L^+(\ell^d) + \delta_L^2 V_L(\sigma_L^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, 2M - \ell^*)).$$

$\ell^d \in [0, \ell^*)$ is winning only in states (R, r) with $r \in [2M - \ell^d, 1] \cup [0, \ell^d]$. For $r \in [2M - \ell^d, 1]$

$$\begin{aligned} V_L(\sigma_L^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, r)) &> \Psi_L^+(\ell^d) + \delta_L^2 V_L(\sigma_L^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, 2M - \ell^d)) \\ &= \Psi_L^+(\ell^d) + \delta_L^2 V_L(\sigma_L^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, 2M - r)). \end{aligned}$$

since

$$\begin{aligned}
V_L(\sigma_{L^*}^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, r)) &= u_L(\ell^*) + \frac{\delta_L}{1 - \delta_L^2} \Psi_L^-(y^1) \\
&> \frac{1}{1 - \delta_L^2} \Psi_L^+(\ell^*) \\
&> \frac{1}{1 - \delta_L^2} \Psi_L^+(\ell^d).
\end{aligned}$$

The first inequality follows from Lemma 1 and the fact that $y^1 > \ell^*$, and the second inequality from Lemma 1 and the fact that $\ell^d < \ell^*$. That a deviation to $\ell^d \in [0, \ell^*)$ in states (R, r) with $r \in [0, \ell^d]$ is not profitable follows from an argument similar to that in Lemma 4. The value of setting $\ell^d \in [y^{i-1}, y^i]$ for $i > 0$ odd if winning at (R, r) is

$$\Psi_L^+(\ell^d) + \delta_L^2 V_L(\sigma_{L^*}^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, 2M - \ell^d)).$$

$\ell^d \in [y^{i-1}, y^i]$ is winning only in states (R, r) with $r \in [2M - \ell^d, 1] \cup [0, \ell^d]$. Consider

$$\begin{aligned}
V_L(\sigma_{L^*}^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, 2M - y^i)) &= u_L(y^i) + \frac{\delta_L}{1 - \delta_L^2} \Psi_L^-(y^{i+1}) \\
&= \frac{1}{1 - \delta_L^2} \Psi_L^+(y^{i-1}) \\
&= u_L(y^{i-1}) + \delta_L u_L(2M - y^{i-1}) + \delta_L^2 V_L(\sigma_{L^*}^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, 2M - y^{i-1})) \\
&\geq \Psi_L^+(\ell^d) + \delta_L^2 V_L(\sigma_{L^*}^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, 2M - \ell^d)),
\end{aligned}$$

where the inequality follows from Lemma 1 and the fact that $\ell^* < y^{i-1} \leq \ell^d$ and the fact that $V_L(\sigma_{L^*}^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, 2M - y^{i-1})) = V_L(\sigma_{L^*}^{\hat{\ell}}, \sigma_R^{\hat{\ell}}; (R, 2M - \ell^i))$. Hence, the value to ℓ^d is weakly smaller than the value following action $y^i = \sigma_L(R, 2M - y^i)$, and hence, by Lemma 5, for all states (R, r) with $r \in [2M - \ell^d, 1]$ deviation to ℓ^d by L cannot be profitable. That a deviation to $\ell^d \in [y^{i-1}, y^i]$ in states (R, r) with $r \in [0, \ell^d]$ is not profitable follows from an argument similar to that in Lemma 4, as does the argument that there is no profitable deviation to $\ell^d \in (M, 1]$.

Arguments very similar to those for L above can determine R 's payoffs under $(\sigma_{L^*}^{\hat{\ell}}, \sigma_R^{\hat{\ell}})$ and verify that it constitutes an equilibrium. Clearly $(\hat{\ell}, 2M - \hat{\ell})$ are robust long-run policy outcomes under $(\sigma_{L^*}^{\hat{\ell}}, \sigma_R^{\hat{\ell}})$ since policy dynamics have $(\hat{\ell}, 2M - \hat{\ell})$ as limit points starting from all more extreme states. □

Let Y be the set of increasing extended real-valued sequences.

Definition 8. Define mapping $B : (\ell^*, M] \rightarrow Y$ such that $B(y)^0 = \ell^*$, $B(y)^1 = y$, for each $i \geq 2$ with i even $B(y)^i$ solves

$$\Psi_L^+(B(y)^{i-2}) - \Psi_L^+(B(y)^{i-1}) = \delta_L [\Psi_L^-(B(y)^i) - \Psi_L^-(B(y)^{i-1})], \quad (18)$$

and for each $i \geq 3$ with i odd, $B(y)^i$ solves

$$\Psi_R^+(B(y)^{i-2}) - \Psi_R^+(B(y)^{i-1}) = \delta_R[\Psi_R^-(B(y)^i) - \Psi_R^-(B(y)^{i-1})], \quad (19)$$

if solutions $B(y)^i \leq M$ exist to (18) and/or (19). If not, set $B(y)^i = \infty$ for all $j \geq i$. Define mapping $\Gamma : (\ell^*, M] \rightarrow \mathbf{R} \cup \{\infty\}$ such that $\Gamma(y) = \lim_{i \rightarrow \infty} B^i(y)$.

Equations (18) and (19) restate the payoff conditions of Lemma 8. Suppose that $\ell^* \geq 2M - r^*$ and that there exists a consistent equilibrium under which $(\hat{\ell}, 2M - \hat{\ell})$ with $\hat{\ell} \in (\ell^*, \ell^{**}]$ is a robust long-run policy outcome. In that case, there exists a convergence path $\{y^i\} \rightarrow \hat{\ell}$ from state (L, ℓ^*) . Suppose that in state (L, ℓ^*) party R selects policy $2M - y$ for $y \in (\ell^*, M]$. The mapping B recovers the full sequence of equilibrium convergence path policies. When no such path exists, we have $B(y)^i = \infty$ for some i . Iteration on B yields a candidate for the sequence posited in Lemma 12, which is acceptable if the limit of $B(y)$, that is $\Gamma(y)$, is contained in $(\ell^*, \ell^{**}]$.

Lemma 13. *Mapping B is such that*

- i. The mapping Γ is well-defined, increasing, strictly increasing on $\{y : \Gamma(y) < \infty\}$, right-continuous on $\{y : \Gamma(y) < \ell^{**}\}$ and left-continuous on $\{y : \Gamma(y) < \infty\}$.*
- ii. For any $\hat{\ell} \in (\ell^*, \ell^{**}]$, there exists y such that $\Gamma(y) = \hat{\ell}$.*
- iii. A strictly increasing sequence $\{y^i\} \rightarrow \hat{\ell}$ with $y^0 = \ell^*$ and y^i, y^{i+1} and y^{i+2} satisfying the conditions of Lemma 8 for all $i \geq 1$.*

Proof: For $y^1 \in (\ell^*, M]$, $\Gamma(y^1)$ is the limit an increasing extended real-valued sequence and hence is well-defined. For the monotonicity of Γ , consider $y^1, \tilde{y}^1 \in (\ell^*, M]$ such that $y^1 < \tilde{y}^1$, along with induced sequences $\{B(y^1)^i\} = \{y^i\}$ and $\{B(\tilde{y}^1)^i\} = \{\tilde{y}^i\}$. First show that for $i \geq 1$, whenever $\infty > \tilde{y}^{i-1} \geq y^{i-1}$, $\infty > \tilde{y}^i > y^i$, $\tilde{y}^i - \tilde{y}^{i-1} > y^i - y^{i-1}$, and $y^{i+1}, \tilde{y}^{i+1} < \infty$, it is the case that $\tilde{y}^{i+1} - \tilde{y}^i > y^{i+1} - y^i$ and $\tilde{y}^{i+1} > y^{i+1}$. Suppose $\tilde{y}^{i-1} - \epsilon = y^{i-1}$, where $\epsilon \geq 0$. Hence

$$\begin{aligned} \Psi_L^+(\tilde{y}^{i-1} - \epsilon) - \Psi_L^+(\tilde{y}^i - \epsilon) &= \delta_L[\Psi_L^-(y^{i+1}) - \Psi_L^-(\tilde{y}^i - \epsilon)] \\ &> \Psi_L^+(y^{i+1}) - \Psi_L^+(y^i) - \delta_L[\Psi_L^-(y^{i+1}) - \Psi_L^-(y^i)] \\ &= 0, \end{aligned}$$

where the inequality follows by Lemma 1 since $\tilde{y}^i - y^i > \epsilon$. Define $\bar{y}^{i+1}$ such that

$$\Psi_L^+(\bar{y}^{i-1} - \epsilon) - \Psi_L^+(\bar{y}^i - \epsilon) - \delta_L[\Psi_L^-(\bar{y}^{i+1}) - \Psi_L^-(\bar{y}^i - \epsilon)] = 0.$$

It must be that $\bar{y}^{i+1} > y^{i+1}$. By Lemma 1, it is also the case that

$$\begin{aligned} \Psi_L^+(\bar{y}^{i-1}) - \Psi_L^+(\bar{y}^i) &= \delta_L[\Psi_L^-(\bar{y}^{i+1} + \epsilon) - \Psi_L^-(\bar{y}^i)] \\ &> \Psi_L^+(\bar{y}^{i-1} - \epsilon) - \Psi_L^+(\bar{y}^i - \epsilon) - \delta_L[\Psi_L^-(\bar{y}^{i+1}) - \Psi_L^-(\bar{y}^i - \epsilon)] \\ &= 0, \end{aligned}$$

and hence $\tilde{y}^{i+1} > \bar{y}^{i+1} + \epsilon > y^{i+1}$ and $\tilde{y}^{i+1} - \tilde{y}^i > \bar{y}^{i+1} - \bar{y}^i - \epsilon > y^{i+1} - y^i$. By induction, if $y^1, \tilde{y}^1 \in \{y : \Gamma(y) < \infty\}$, this implies that for each $i \geq 1$, $\tilde{y}^i > y^i$, and

$$\begin{aligned}\Gamma(\tilde{y}^1) &= \lim_{i \rightarrow \infty} \tilde{y}^i \\ &> \lim_{i \rightarrow \infty} y^i \\ &= \Gamma(y^1).\end{aligned}$$

The above argument also shows that if $y^1 < \tilde{y}^1$, then $y^i < \tilde{y}^i$ for all i such that $\tilde{y}^i < \infty$, and hence that $\Gamma(y^1) \leq \Gamma(\tilde{y}^1)$.

Suppose Γ is not right-continuous at y^1 , and that $\Gamma(y^1) < \ell^{**}$. Then there exists $\epsilon > 0$ such that for any $\delta > 0$, $\Gamma(y^1 + \delta) - \Gamma(y^1) > \epsilon$. Take $\bar{\epsilon} \in (0, \min\{\epsilon, \ell^{**} - \Gamma(y^1)\})$. Hence $\Gamma(y^1) + \bar{\epsilon} < \ell^{**}$. Consider $\tilde{y}^1 \in (y^1, y^1 + \delta)$ and associated sequence $\{\tilde{y}^i\}$. Since $\Gamma(y^1) + \bar{\epsilon} < \ell^{**}$, by part ii of Lemma 11 there exist α_L and α_R with $\alpha_L + \alpha_R > 1$ such that for any $\{\bar{y}^i\} \rightarrow \Gamma(\tilde{y}^1)$ with $\Gamma(\bar{y}^1) \leq \Gamma(y^1) + \bar{\epsilon}$, $\bar{y}^{i+1} - \bar{y}^i < \frac{\alpha_L}{1-\alpha_L}(\bar{y}^i - \bar{y}^{i-1})$, $\bar{y}^i - \bar{y}^{i-1} < \frac{\alpha_R}{1-\alpha_R}(\bar{y}^{i-1} - \bar{y}^{i-2})$ and

$$\lim_{i \rightarrow \infty} \bar{y}^i < \bar{y}^0 + (\bar{y}^1 - \bar{y}^0) \frac{\frac{\alpha_L}{1-\alpha_L}(1 + \frac{\alpha_R}{1-\alpha_R})}{1 - \frac{\alpha_L}{1-\alpha_L} \frac{\alpha_R}{1-\alpha_R}}.$$

Conversely, if $\bar{y}^0 + (\bar{y}^1 - \bar{y}^0) \frac{\frac{\alpha_L}{1-\alpha_L}(1 + \frac{\alpha_R}{1-\alpha_R})}{1 - \frac{\alpha_L}{1-\alpha_L} \frac{\alpha_R}{1-\alpha_R}} \leq \Gamma(y^1) + \bar{\epsilon}$, then it must be that $\Gamma(\tilde{y}^1) < \Gamma(y^1) + \bar{\epsilon}$.

Since $\{y^i\} \rightarrow \Gamma(y^1)$, there exists $n \in \mathbf{N}$ such that

$$y^i + (y^{i+1} - y^i) \frac{\frac{\alpha_L}{1-\alpha_L}(1 + \frac{\alpha_R}{1-\alpha_R})}{1 - \frac{\alpha_L}{1-\alpha_L} \frac{\alpha_R}{1-\alpha_R}} < \Gamma(y^1) + \frac{\bar{\epsilon}}{2}$$

for all $i \geq n$. Fix $j \geq n$. Since for all $i \geq 1$, $\tilde{y}^{i+1}$ is a continuous function of $\tilde{y}^i$ and $\tilde{y}^{i-1}$, $\tilde{y}^1$ can be found such that $\tilde{y}^j - y^j < \frac{\bar{\epsilon}}{4}$ and $(\tilde{y}^{j+1} - \tilde{y}^j) - (y^{j+1} - y^j) < \frac{\bar{\epsilon}}{4} \frac{1 - \frac{\alpha_L}{1-\alpha_L} \frac{\alpha_R}{1-\alpha_R}}{\frac{\alpha_L}{1-\alpha_L}(1 + \frac{\alpha_R}{1-\alpha_R})}$. Then it follows that

$$\begin{aligned}\tilde{y}^j + (\tilde{y}^{j+1} - \tilde{y}^j) \frac{\frac{\alpha_L}{1-\alpha_L}(1 + \frac{\alpha_R}{1-\alpha_R})}{1 - \frac{\alpha_L}{1-\alpha_L} \frac{\alpha_R}{1-\alpha_R}} &< y^j + \frac{\bar{\epsilon}}{4} + (y^{j+1} - y^j) \frac{\frac{\alpha_L}{1-\alpha_L}(1 + \frac{\alpha_R}{1-\alpha_R})}{1 - \frac{\alpha_L}{1-\alpha_L} \frac{\alpha_R}{1-\alpha_R}} + \frac{\bar{\epsilon}}{4} \\ &< \Gamma(y^1) + \bar{\epsilon}.\end{aligned}$$

Hence $\Gamma(\tilde{y}^1)$ is such that $\Gamma(\tilde{y}^1) < \Gamma(y^1) + \bar{\epsilon}$, a contradiction.

Suppose Γ is not left-continuous at y^1 , and that $\Gamma(y^1) < \infty$. Then there exists $\epsilon > 0$ such that for any $\delta > 0$, $\Gamma(y^1) - \Gamma(y^1 - \delta) > \epsilon$. Take $j \in \mathbf{N}$ such that $y^j > \Gamma(y^1) - \epsilon + \eta$ for $\eta \in (0, \epsilon)$. Fix $\tilde{y}^1$ such that $y^j - \tilde{y}^j < \eta$. Hence $\tilde{y}^j > y^j - \eta > \Gamma(y^1) - \epsilon$, and hence $\Gamma(\tilde{y}^1) > \Gamma(y^1) - \epsilon$, since $\{\tilde{y}^i\}$ is increasing, a contradiction.

The set $\{y : \Gamma(y) < \ell^{**}\}$ is nonempty since $\lim_{y \rightarrow \ell^*} \Gamma(y) = \ell^*$, and hence by continuity of Γ on $\{y : \Gamma(y) < \ell^{**}\}$, for each ℓ with $\ell < \ell^{**}$, there exists y such that $\Gamma(y) = \ell$. Finally, since Γ is left-continuous on $\{y : \Gamma(y) < \infty\}$, there exist a y such that $\Gamma(y) = \ell^{**}$.

□

A.5 Forward-looking Voters

Proof of Proposition 5: Consider consistent equilibrium convergence path $\{y^i\}$ with associated consistent equilibrium strategies (σ_L, σ_R) . Assume for now that on convergence paths, the median voter votes according to σ_M^{my} . To construct strategies (σ'_L, σ'_R) in the game with forward-looking voters, the profile (σ_L, σ_R) needs to be modified in two ways. First, consider policy y^i such that $\sigma_L(R, 2M - y^i) = y^{i+1}$. For $x \in [y^i, y^{i+1})$, define $z^{i+1}(x) \in [y^i, x)$ such that

i. If

$$u_M(x) - u_M(y^i) > \delta_M \left[V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (R, 2M - y^i)) - \frac{1}{1 - \delta_M} u_M(x) \right],$$

then $z^{i+1}(x)$ solves

$$u_M(x) - u_M(z^{i+1}(x)) = \delta_M \left[V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (R, 2M - y^i)) - \frac{1}{1 - \delta_M} u_M(x) \right].$$

ii. If

$$u_M(x) - u_M(y^i) \leq \delta_M \left[V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (R, 2M - y^i)) - \frac{1}{1 - \delta_M} u_M(x) \right],$$

then $z^{i+1}(x) = y^i$.

That is, R commits to $2M - z^{i+1}(x)$ as ‘punishment’ for L being in power with policy x as opposed to y^{i+1} and $z^{i+1}(x)$ is the most extreme such punishment that the median voter supports. For y^i such that $\sigma_R(L, y^i) = 2M - y^{i+1}$ and for $x \in (2M - y^{i+1}, 2M - y^i]$, $z^{i+1}(x) \in [y^i, 2M - x)$ can be defined symmetrically.

Second, given some σ_M and $\ell > M$, let $\bar{r}(\ell) > \ell$ be the most extreme commitment by R in state (L, ℓ) that the median voter supports and that R has the incentive to make. If the median voter accepts $\bar{r}(\ell)$, then policy dynamics are ‘freed’ from the policy traps of equilibria with myopic voters and, after at most one period, the equilibrium path rejoins convergence path $\{y^i\}$. For $r < M$, define $\bar{\ell}(r) < r$ symmetrically. Note that, as with the functions $\{z^{i+1}(\cdot)\}$, $\bar{r}(\cdot)$ and $\bar{\ell}(\cdot)$ are determined only by how parties and the median voter evaluate convergence paths under $(\sigma_L, \sigma_R, \sigma_M^{my})$. Now define strategy σ'_R as

$$\sigma'_L(R, r) = \begin{cases} z^{i+1}(r) & \text{if } r \in (2M - y^{i+1}, 2M - y^i] \text{ for } y^i \text{ such that } \sigma_R(L, y^i) = 2M - y^{i+1}, \\ \bar{\ell}(r) & \text{if } r < M \text{ and } u_L(\bar{\ell}(r)) + \delta_L V_L(\sigma_L, \sigma_R; (L, \bar{\ell}(r))) \geq \frac{1}{1 - \delta_L} u_L(r) \\ \sigma_L(R, r) & \text{otherwise.} \end{cases}$$

σ'_R can be defined symmetrically. Let σ_M be a best-response to (σ'_L, σ'_R) in which the median voter supports the opposition party when indifferent. Given the parties’ strategies, the median

voter has no incentive to vote for the incumbent on a convergence path. Hence, given convergence path policy y^i such that $\sigma_L(R, 2M - y^i) = y^{i+1}$, we have that $V_K(\sigma'_L, \sigma'_R, \sigma_M; (R, 2M - y^i)) = V_K(\sigma_L, \sigma_R, \sigma_M^{my}; (R, 2M - y^i))$ for $K \in \{L, R, M\}$. I do not describe the median voter's equilibrium strategy explicitly, but instead show how it responds to parties' deviations from the convergence path $\{y^i\}$ to show that parties have no more incentive to deviate from the convergence path under $(\sigma'_L, \sigma'_R, \sigma_M)$ than under $(\sigma_L, \sigma_R, \sigma_M^{my})$.

Consider state (R, r) with $2M - r \in [y^i, y^{i+1})$ for y^i such that $\sigma_R(L, y^i) = 2M - y^{i+1}$. The median voter votes against $\ell \in [y^i, z^{i+1}(r))$ since the payoff to voting in favour of ℓ is

$$u_M(\ell) + \delta_M V_M(\sigma'_L, \sigma'_R, \sigma_M; (L, y^i)) < u_M(r) + \delta_M u_M(z^{i+1}(r)) + \delta_M^2 V_M(\sigma'_L, \sigma'_R, \sigma_M; (L, y^i)),$$

by the definition of $z^{i+1}(r)$, where the right-hand side is the payoff to voting in favour of r . The median voter votes against $\ell \in [y^{i-1}, y^i)$ since the payoff to voting in favour of ℓ is

$$\begin{aligned} u_M(\ell) + \delta_M u_M(z^i(\ell)) + \delta_M^2 V_M(\sigma'_L, \sigma'_R, \sigma_M; (R, 2M - y^{i-1})) \\ < u_M(r) + \delta_M u_M(z^{i+1}(r)) + \delta_M^2 V_M(\sigma'_L, \sigma'_R, \sigma_M; (L, y^i)), \end{aligned}$$

since $|M - \ell| > |M - r|$, $|M - z^i(\ell)| > |M - z^{i+1}(r)|$ and $V_M(\sigma'_L, \sigma'_R, \sigma_M; (R, 2M - y^{i-1})) < V_M(\sigma'_L, \sigma'_R, \sigma_M; (L, y^i))$. Similarly, the median voter votes against $\ell \in [y^{k-1}, y^k)$ for y^k such that $\sigma_L(R, 2M - y^{k-1}) = y^k$ and $k \leq i - 2$, and against $\ell \in [y^{k-1}, y^k)$ for y^k such that $\sigma_R(L, y^{k-1}) = 2M - y^k$ and $k \leq i - 1$. That is, in state (R, r) , the median voter rejects all policies $\ell \in [0, z^{i+1}(r))$. It may or may not vote for policies $\ell \in (z^{i+1}(r), 1]$. A similar argument shows that in state (R, r) with $2M - r \in [y^i, y^{i+1})$ for y^i such that $\sigma_L(R, 2M - y^i) = y^{i+1}$, the median voter rejects any $\ell \in [0, r]$ and may or may not support $\ell \in (r, 1]$, but always supports $\ell = y^{i+1}$.

Now consider parties' incentives. First, whenever a party's equilibrium policy is being accepted, it never gains by committing to policies that are sure to be rejected, since it faces the same choice in the next election. Consider again state (R, r) with $2M - r \in [y^i, y^{i+1})$ for y^i such that $\sigma_R(L, y^i) = 2M - y^{i+1}$. The payoff to party L from policy $\ell \in [z^{i+1}(r), y^{i+1}]$ that is accepted by the median voter is

$$u_L(\ell) + \delta_L u_L(2M - y^{i+1}) + \delta_L^2 V_L(\sigma'_L, \sigma'_R, \sigma_M; (R, 2M - y^{i+1})),$$

which is decreasing in $\ell \in [y^i, y^{i+1})$. From above, policies $\ell \in [0, z^{i+1}(r))$ cannot be profitably proposed since they are rejected by the median voter, while policies in $(y^{i+1}, M]$, if accepted, yield to party L at most the payoff it obtains from such deviations under $(\sigma_L, \sigma_R, \sigma_M^{my})$. Hence committing to $z^{i+1}(r)$ is optimal for party L .

Now consider policy y^i such that $\sigma_L(R, 2M - y^i) = y^{i+1}$ and state (R, r) with $2M - r \in [y^i, y^{i+1})$. The payoff from $\ell \in [2M - r, y^{i+1})$, if accepted by the median voter, is given by

$$\begin{aligned} u_L(\ell) + \delta_L u_L(2M - z^{i+1}(\ell)) + \delta_L^2 V_L(\sigma'_L, \sigma'_R, \sigma_M; (R, 2M - y^i)) \\ \leq u_L(\ell) + \delta_L u_L(2M - \ell) + \delta_L^2 V_L(\sigma'_L, \sigma'_R, \sigma_M; (R, 2M - y^i)) \\ < V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (R, 2M - y^i)). \end{aligned}$$

The first inequality follows from $z^{i+1}(\ell) \leq \ell$ and the second since $V_L(\sigma'_L, \sigma'_R, \sigma_M; (R, 2M - y^i)) > \frac{1}{1-\delta_L^2} \Psi_L^+(\ell)$. This shows that y^{i+1} is L 's preferred winning policy in $[y^i, y^{i+1})$ given $(\sigma'_L, \sigma'_R, \sigma_M)$. As the median voter rejects any policy $\ell \in [0, 2M - r)$, L cannot profitably deviate to such policies. Finally, deviations to any policies $\ell \in (y^{i+1}, M]$ are never profitable since even if they are accepted by the median voter, L 's payoffs are no higher than under $(\sigma_L, \sigma_R, \sigma_M^{my})$.

It remains to deal with states (R, r) with $r < M$. By construction, in these states σ'_L is optimal. It needs to be shown that in states (R, r) with $r \geq M$, party L does not want to deviate to some $\ell^d > M$. Consider state (R, r) with $r > M$, and suppose party L deviates to $\ell^d > M$ such that $\sigma'_R(L, \ell^d) = \bar{r}(\ell^d)$ and take $\{y^i\}$ to be the convergence path from $(R, \bar{r}(\ell^d))$. It must be that $y^1 \geq 2M - \bar{r}(\ell^d)$. The payoff to party L from ℓ^d is given by

$$\begin{aligned} u_L(\ell^d) + \delta_L u_L(\bar{r}(\ell^d)) + \sum_{i=1}^{\infty} \delta^{2i+1} [u_L(y^i) + \delta_L u_L(2M - y^{i+1})] &< u_L(\ell^d) + \frac{\delta_L}{1 - \delta_L} u_L(M) \\ &< \frac{1}{1 - \delta_L} u_L(M). \end{aligned}$$

The first inequality follows by Lemma 1 and the second since $\ell^d > M$. On the equilibrium path, $V_L(\sigma_L, \sigma_R; (R, r)) \geq \frac{1}{1-\delta_L} u_L(M)$, and hence deviation to ℓ^d is not profitable for L . $\square$

A.6 Legislative Bargaining

Proof of Proposition 6: Consider consistent proposal strategies (σ_L, σ_R) that generate convergence path $\{y^i\} \rightarrow \hat{\ell}$ when the median legislator is decisive and $\sigma_M = \sigma_M^{my}$. It will be shown that σ_M^{my} is indeed a best response for the median legislator. It is straightforward to establish results equivalent to Lemma 7 that characterises consistent proposal strategies on convergence paths.

Consider a convergence path $\{y^i\} \rightarrow \hat{\ell}$ with policy y^i such that $\sigma_L(L, 2M - y^i) = y^{i+1}$. Since each legislator is recognised with equal probability in each period, legislator L 's equilibrium payoff is given by

$$\begin{aligned} V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (L, 2M - y^i)) &= u_L(y^{i+1}) + \frac{1}{2} \delta_L V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (L, y^{i+1})) \\ &\quad + \frac{1}{2} \delta_L V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (R, y^{i+1})) \\ &= \frac{2}{2 - \delta_L} \left[u_L(y^{i+1}) + \frac{1}{2} \delta_L V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (R, y^{i+1})) \right], \quad (20) \end{aligned}$$

where the second equality is due to consistent proposal strategies. A lower bound on $V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (L, 2M - y^i))$ can be determined as in Section 4.3 by considering a deviation

to y^i by L in state $(L, 2M - y^i)$. Hence

$$\begin{aligned} V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (L, 2M - y^i)) &\geq u_L(y^i) + \frac{1}{2}\delta_L V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (L, 2M - y^i)) \\ &\quad + \frac{1}{2}\delta_L V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (R, y^i)). \end{aligned} \quad (21)$$

By convergence and consistent strategies, $\sigma_R(R, y^i) = \sigma_R(R, 2M - y^i) = 2M - y^i$, and hence, as for (20) above,

$$V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (R, y^i)) = \frac{2}{2 - \delta_L} \left[u_L(2M - y^i) + \frac{1}{2}\delta_L V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (L, 2M - y^i)) \right]. \quad (22)$$

Under consistent strategies, an upper bound on $V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (L, 2M - y^i))$ can be obtained as in Section A.3 by considering a deviation to y^{i+1} in state $(R, 2M - y^i + \epsilon)$ for small ϵ . That is

$$\begin{aligned} V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (L, 2M - y^i)) &\geq u_L(y^i) + \frac{1}{2}\delta_L V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (L, 2M - y^i)) \\ &\quad + \frac{1}{2}\delta_L V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (R, y^i)). \end{aligned} \quad (23)$$

Finally, (21), (23) and (22) yield

$$V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (L, 2M - y^i)) = \frac{2 - \delta_L}{2(1 - \delta_L)} \left[u_L(y^i) + \frac{\delta_L}{2 - \delta_L} u_L(2M - y^i) \right]. \quad (24)$$

This is the equivalent of (8) which states that L 's equilibrium payoff at $(R, 2M - y^i)$ is the payoff to alternation at $(y^i, 2M - y^i)$. Expression (24) incorporates the fact that the future sequence of proposers is random and that convergence is staggered. A calculation like the one in (22) yields $V_L(\sigma_L, \sigma_R, \sigma_M^{my}; (R, y^{i+1}))$, and (20) can be rewritten, after substituting (24), as

$$\begin{aligned} (2 - \delta_L) [u_L(y^i) - u_L(y^{i+1})] &+ \frac{\delta_L^2}{2 - \delta_L} [u_L(y^{i+1}) - u_L(y^{i+2})] \\ &= \delta_L [u_L(2M - y^{i+2}) - u_L(2M - y^i)]. \end{aligned} \quad (25)$$

Equation (25) is the equivalent of (9), the second-order differential equation that determines consistent equilibrium convergence path policies, in the legislative bargaining model. Conditions for existence of convergence paths in this model would hinge on the properties of the payoffs of legislators L and R relative to (25). However, for the purposes of Proposition 6, all that is required is that (25) must hold along any convergence path in consistent proposal strategies.

As in A.4, a bound on the moderation of convergence outcomes can be derived by constructing 'compromise' functions α_L and α_R . An argument as in A.4 shows that given some $y < M$, $\alpha_L(y)$ can be defined as

$$\frac{u'_L(y)}{u'_L(2M - y)} = \frac{\delta_L}{\alpha_L(y)(2 - \delta_L) + (1 - \alpha_L(y))\frac{\delta_L^2}{2 - \delta_L}}.$$

In particular, $\alpha_L(M) = \frac{\delta_L}{2} < \frac{1}{2}$, and a similar argument shows that $\alpha_R(M) < \frac{1}{2}$. Hence, as in Section A.4, as convergence paths approach the median, both legislators require that their opponent's next moderate move be larger than their own current moderate move, which contradicts convergence.

I have assumed that median voter behaves myopically. In fact, it can be shown that this voting strategy is optimal. Consider policy y^i such that $\sigma_L(L, 2M - y^i) = y^{i+1}$. Suppose that in state $(L, 2M - y^i)$ legislator L proposes $z \in [y^i, y^{i+1})$. If the median voter votes in favour of z its payoff is given by

$$\begin{aligned} u_M(z) &+ \frac{1}{2}\delta_M V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (L, z)) + \frac{1}{2}\delta_M V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (R, z)) \\ &> u_M(2M - y^i) + \frac{1}{2}\delta_M V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (L, 2M - y^i)) + \frac{1}{2}\delta_M V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (R, 2M - y^i)), \end{aligned}$$

where the right-hand side is the payoff to supporting the status quo. This follows since $u_M(z) > u_M(2M - y^i)$, $V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (L, z)) = V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (L, 2M - y^i))$ since $\sigma_L(L, z) = \sigma_L(2M - y^i) = y^{i+1}$ and $V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (R, z)) > V_M(\sigma_L, \sigma_R, \sigma_M^{my}; (R, 2M - y^i))$ since $\sigma_R(R, \ell) = 2M - \ell$ for $\ell \in [y^i, y^{i+1})$. Similar arguments show that the median legislator accepts any policy $z \in [y^{i+1}, 2M - y^{i+1}]$ and rejects any policy $z \in [y^i, 2M - y^i]$. Furthermore, these arguments do not depend on which legislator makes the proposal, since future periods' draws of proposers are not affected by the identity of the legislator responsible for the status quo policy. □