

Majority Bargaining and Reputation: The Symmetric Case

Zizhen Ma
Dept. of Economics
University of Rochester

February 28, 2018

Abstract

This paper analyzes the interaction of non-unanimity and reputation in a simple environment. Three agents bargain over the division of one dollar with majority rule and a Baron-Ferejohn protocol with uniform recognition probabilities. Each agent could be a semi-rational obstinate type committed to claim a certain share of the dollar. In sharp contrast to bilateral bargaining with reputational concerns, assuming common conflicting claims and a common discount factor, I show that when rational types are sufficiently patient, there is a perfect Bayesian equilibrium in which the bargaining process is asymptotically efficient in the sense that it reaches a potential agreement in finitely many periods with probability one. Moreover, in this equilibrium, the rational type with the weakest reputation of obstinacy obtains the largest share of the dollar.

1 Introduction

Agents in bargaining may behave obstinately as intrinsic behavioral pattern or out of strategic concern of reputation formation. For bilateral bargaining, it is well-known that with reputational concerns about obstinacy: agents are necessarily engaged in a war of attrition which results in significant efficiency loss; moreover, the agent with stronger reputation of obstinacy wins the larger share of the surplus. However, efficiency and distribution implications of such reputational concerns remain unknown for bargaining environment in which more than two agents are involved and unanimity is not required for settlement, while non-unanimous bargaining is common and important for decision making in legislatures, international relations, and various committees (in government, company or academia).

This paper studies surplus division with reputational concerns in majority bargaining: three agents bargain over the division of one dollar with majority rule and the infinite-horizon Baron-Ferejohn protocol with uniform recognition probabilities, and each agent could be a semi-rational obstinate type committed to claim a certain share of the dollar. Assuming common conflicting claims and a common discount factor, I show by construction that when rational types are sufficiently patient or when bargaining interaction is sufficiently frequent, there is a perfect Bayesian equilibrium in which the bargaining process is asymptotically efficient in the sense that it reaches a potential agreement in finitely many periods with probability one. Intuitively, threat of exclusion in non-unanimous bargaining eliminates

incentives of reputation formation and significantly improves efficiency. Moreover, when the prior probability of semi-rationality is small, in the constructed equilibrium the rational type with the weakest reputation of obstinacy wins the largest share of surplus as he can guarantee inclusion in all winning coalitions formed on equilibrium path.

These two features stand in sharp contrast to their counterparts in unanimous bargaining as mentioned earlier. The seminal works on bargaining and reputation by Kambe (1999) and Abreu and Gul (2000) show that in bilateral bargaining with two-sided reputational concerns, agents are necessarily engaged in a war of attrition which results in significant efficiency loss even when bargaining interaction is frequent. When irrational types have common conflicting claims and the prior probability of irrationality is small, their models predict that the rational type with the stronger reputation wins the larger share. My equilibrium construction adopts technique developed in the literature of bilateral bargaining with non-reputational incomplete information, and here I refer to Ausubel et al. (2002) for an excellent survey.

The literature on non-unanimous bargaining with complete information has been well-developed since Baron and Ferejohn (1989) defines the later widely used Baron-Ferejohn protocols. Banks and Duggan (2000, 2006) show existence of stationary equilibria and core convergence for general sets of alternatives. Eraslan (2002) and Eraslan and McLennan (2012) show that stationary equilibrium payoffs of the divide the dollar game is unique. Threat of exclusion in majority bargaining has received attention and yield interesting insights in deviations from the standard environment. Eraslan and Merlo (2002) show that when the value of one dollar varies stochastically over time, threat of exclusion may lead to inefficient early agreement. Ali (2006) shows that when agents disagree over subjective assessment of recognition probabilities, threat of exclusion may improve efficiency instead.

Comparatively, non-unanimous bargaining with incomplete information is less explored. Tsai and Yang (2010) study a three-period three-agent divide-the-dollar model in which each agent has private information about his patience. They show that delay may happen due to signaling incentive, and information revelation is determined by tradeoff between terms of agreement and probability of inclusion: an impatient type obtains less from agreement but is included with high probability.¹ Chen and Eraslan (2014) studies two-period three-agent legislative bargaining over ideological and distributive issues. Each agent has private information about his ideological intensity. They show that allowing all agents to send cheap-talk messages to influence the proposer may crowd out information in equilibrium and thus make the proposer worse off. Chen (2017) studies a two-period legislative bargaining model with unidimensional alternatives and a fixed agenda setter. Each voter has private information about his ideal point. He shows that the agenda setter can extract information about voters' preferences by a tentative proposal in the first period.

The remainder of the paper is organized as follows. Section 2 defines the model of majority bargaining with reputation. Section 3 presents the main results on existence, efficiency and surplus distribution. Section 4 presents the key steps of equilibrium construction and the proofs of the main results. Section 5 concludes and discusses future research directions. Section 6 as Appendix contains detailed proofs of some propositions.

¹This paper also highlights this tradeoff, but the reason why the rational type with the weakest reputation of commitment is most advantageous is more subtle. See Section 3 for a detailed discussion.

2 Model

In this section, the model of majority bargaining with reputation is defined.

A set of agents $N = \{1, 2, 3\}$ bargain over the division of one dollar with majority rule and the infinite-horizon Baron-Ferejohn protocol with uniform recognition probabilities. Bargaining starts at period 0. The timing of interaction at each period t prior to agreement is as follows: (1) agent $i \in N$ is recognized as the proposer with probability $\frac{1}{3}$; then (2) agent i proposes $x \in X = \{x \in \mathbb{R}_+^N \mid \sum_{i \in N} x_i = 1\}$; then (3) every agent $j \in N$ votes simultaneously to either accept or reject the proposal. If x is accepted by at least two agents at period t , bargaining ends with outcome (x, t) ; otherwise, the game continues to period $t + 1$, and the process is repeated.

Each agent $i \in N$ can be either rational or semi-rational. The rational type of agent i evaluates a bargaining outcome (x, t) according to $u_i(x, t) = \delta^t x_i$, with $\delta \in (0, 1)$ as the common discount factor. The semi-rational type of agent i only makes proposals in $X_c^i = \{x \in X \mid x_i = c\}$ with $c \in (\frac{1}{2}, 1)$ representing common conflicting claims, and he accepts a proposal x if and only if $x_i \geq c$. Note that X_c^i is not a singleton, and I postpone describing the proposing behavior of a semi-rational type until I formulate the equilibrium concept. Nature draws the types of agents independently according to $\mu \in [0, 1]^N$, with μ_i being the probability that agent i is semi-rational. The knowledge about whether agent i is rational or semi-rational is private to agent i , and all other features of the game is common knowledge. When $\mu_i = 0$, the rational type of agent i is referred to as revealed rational; when $\mu_i = 1$, he is referred to as concealed rational.

I look for a perfect Bayesian equilibrium in which the strategies of rational types are stationary and symmetric in pre-proposal belief of each period.² Stationarity in pre-proposal belief of each period states that the proposal strategy of the rational type of agent i can be represented by a mapping $\pi_i : [0, 1]^N \rightarrow \Delta(X)$, and his voting strategy can be represented by a mapping $\rho_i : N \times X \times [0, 1]^N \rightarrow [0, 1]$ with $\rho_i(\cdot) = 1$ standing for pure rejection. Symmetry in beliefs states that if μ and μ' is related by a permutation pm of N such that for all $i \in N$, $\mu'_i = \mu_{pm(i)}$, then for all $j, j' \in N$, $\pi_j(\mu') = \pi_{pm(j)}(\mu)$ and $\rho_j(j', x, \mu') = \rho_{pm(j)}(pm(j'), x, \mu)$. Stationarity and symmetry in pre-proposal belief at each period are also imposed for semi-rational types, i.e., the proposing behavior of semi-rational types can be represented by a symmetric mapping $\chi : [0, 1]^N \rightarrow \prod_{i \in N} \Delta(X_{\alpha_i})$. Given a stationary profile $\sigma = (\pi, \rho, \chi)$, pre-recognition payoffs and post-recognition payoffs for rational types can be recursively defined. To perform this task, some more formality is developed first. Let $Y \in \{0, 1\}^N$ be the set of voting outcome such that for each $y \in Y$, $y_i = 1$ stands for acceptance by agent i and let $Y_r = \{y \in Y \mid \sum_{i \in N} y_i \leq 1\}$ be the set of voting outcomes associated with rejection. Note that for each $\mu \in [0, 1]^N$, σ defines a probability measure $\tilde{\sigma}(\mu)$ on $N \times X \times Y$, a typical element (i, x, y) of which records the identity of proposer, the proposal and the voting outcome. For each $(i, x, y) \in N \times X \times Y$, let $\Lambda(\mu; \sigma)(i, x, y) \in [0, 1]^N$ be the posterior defined by Bayesian updating according to σ . Then given μ and σ , the pre-recognition payoffs for rational types can be expressed as

$$v(\mu; \sigma) = E_{\tilde{\sigma}(\mu)} [x(1 - \mathbf{1}_{\tilde{y} \in Y_r}) + \mathbf{1}_{\tilde{y} \in Y_r} \delta v(\Lambda(\mu; \sigma)(\tilde{i}, \tilde{x}, \tilde{y}); \sigma)]$$

and the post-recognition payoffs for rational types when agent j proposes can be expressed

²Stationarity in pre-proposal belief permits the flexibility of having voting decisions depend on both pre-proposal belief and post-proposal belief within each period and significantly simplifies equilibrium construction.

as

$$w(j, \mu; \sigma) = E_{\tilde{\sigma}(\mu)(\cdot|j)} [x(1 - \mathbf{1}_{\tilde{y} \in Y_r}) + \mathbf{1}_{\tilde{y} \in Y_r} \delta v(\Lambda(\mu; \sigma)(j, \tilde{x}, \tilde{y}); \sigma)].$$

Now σ is a stationary perfect Bayesian equilibrium if for all $\mu \in [0, 1]^N$ and for all $i, j \in N$,
(1)

$$w_i(i, \mu; \sigma) = \max_{x \in X} E_{\tilde{\sigma}(\mu)(\cdot|i, x)} [x(1 - \mathbf{1}_{\tilde{y} \in Y_r}) + \mathbf{1}_{\tilde{y} \in Y_r} \delta v(\Lambda(\mu; \sigma)(i, x, \tilde{y}); \sigma)],$$

where $\tilde{\sigma}(\mu)(\cdot|i, x)$ is the probability measure over Y defined by ρ , this condition simply requires that the proposals made by the rational type of agent i is optimal in response to σ ; (2) $\rho_i(j, \mu, x) = 0$, if

$$x_i > \delta E_{\tilde{\sigma}(\mu)(\cdot|j, x, y_i=0, \tilde{y} \in Y_r)} [v(\Lambda(\mu; \sigma)(j, x, \tilde{y}); \sigma)],$$

and $\rho_i(j, \mu, x) = 1$, if

$$x_i < \delta E_{\tilde{\sigma}(\mu)(\cdot|j, x, y_i=0, \tilde{y} \in Y_r)} [v(\Lambda(\mu; \sigma)(j, x, \tilde{y}); \sigma)],$$

where $\tilde{\sigma}(\mu)(\cdot|j, x, y_i=0, \tilde{y} \in Y_r)$ is the probability measure over voting outcomes defined by ρ conditional on rejection by i and rejection of x , this condition is an adaptation of *stage dominance* to our environment; (3) for all $x \in \text{supp}(\chi_i)$,

$$E_{\tilde{\sigma}(\mu)(\cdot|i, x)} [(1 - \mathbf{1}_{\tilde{y} \in Y_r})] = \max_{x' \in X_{c_i}^i} E_{\tilde{\sigma}(\mu)(\cdot|i, x')} [(1 - \mathbf{1}_{\tilde{y} \in Y_r})],$$

this condition corresponds to the term "semi-rational" and states that a semi-rational type maximizes the chance of fulfilling his claim.

3 Main theorems

In this section, the two main results are presented. The analysis focuses on environments with sufficiently patient rational types, i.e., when δ is close to 1. The model with sufficiently patient rational types is not only more tractable, but also highlights the interaction between non-unanimity and reputation. Note that the results for patient rational types can also be interpreted as results for frequent bargaining interaction. Let Δ be the time elapse between a rejection and the nearest upcoming recognition, let $r > 0$ be the common discount rate for rational types, and let $\delta = e^{-r\Delta}$; then $\delta \rightarrow 1$ is equivalent to $\Delta \rightarrow 0$.

The first main result states that when rational types are sufficiently patient, there is an equilibrium in which the bargaining process would reach a potential agreement in finitely many periods with probability one. It implies that efficiency loss vanishes as rational types become patient or as bargaining interaction becomes frequent. Thus the bargaining process exhibits a Coasian feature.

Theorem 1 Fix $c \in (\frac{1}{2}, 1)$ and $\mu \in [0, 1]^N$. There is $\delta^*(c, \mu) \in (0, 1)$ such that for all $\delta > \delta^*(c, \mu)$, there is a symmetric stationary perfect Bayesian equilibrium. In this equilibrium, for each $T \geq 4$, the probability that bargaining ends by period T is at least $(1 - \prod_{i \in N} \mu_i) \left[1 - \left(\frac{1}{3}\right)^{T-4}\right]$.

Note that there is a potential agreement if and only if at least one agent is rational, then $1 - \prod_{i \in N} \mu_i$ is the prior probability that there is a potential agreement. And given that there is a potential agreement, Theorem 1 states that bargaining ends in agreement in finitely many period with probability one. Let's consider three simple cases as examples. (1) Suppose $\mu = (1, 1, 1)$, then there is no potential agreement, and indeed the probability that bargaining ends by period T is $0 = (1 - 1) \left[1 - \left(\frac{1}{3}\right)^{T-4}\right]$. (2) Suppose $\mu = (0, 0, 1)$, then there is an equilibrium in which the revealed rational agent 1 and agent 2 seek immediate agreement with each other, while the proposals by agent 3 is always rejected. An agreement is reached by period T if and only if at least one of agent 1 and agent 2 is recognized by period T . And thus the probability that bargaining ends by period T is $1 - \left(\frac{1}{3}\right)^{T+1} \geq (1 - 0) \left[1 - \left(\frac{1}{3}\right)^{T-4}\right]$. (3) Suppose $\mu \in (0, 1) \times \{1\} \times \{1\}$, then bargaining ends immediately if agent i is rational, while there is no potential agreement if agent i is semi-rational. And thus the probability that bargaining ends by period T is $1 - \mu_1 \geq (1 - \mu_1) \left[1 - \left(\frac{1}{3}\right)^{T-4}\right]$.

The main force at work is that if one agent establishes himself as semi-rational, he may be excluded from winning coalitions that the other agents planning to form (as illustrated in the second example above). As semi-rationality may be a curse, when a rational type is proposing, he can take advantage of a threat of denouement of semi-rationality and seek agreement with an aggressive proposal by another rational type, which provides a bonus to revealing rationality for a proposer. Then a rational type in the game may reveal rationality early either actively as a proposer or passively as screened by the others. In contrast, as shown in Abreu and Gul (2000), a bilateral bargaining with two-sided reputation formation and conflicting claims entails slow revelation of rationality and non-trivial efficiency loss even when rational types are patient or bargaining interaction is frequent. Theorem 1 is shown by explicit construction in Section 4.

Fix $c \in (\frac{1}{2}, 1)$ and $\mu \in [0, 1]^N$. For each $\delta > \delta(c, \mu)$, the constructed equilibrium in Theorem 1 defines an equilibrium payoff vector $v^\delta(c, \mu)$ for rational types. Following from the construction,

$$v^*(c, \mu) = \lim_{\delta \rightarrow 1} v^\delta(c, \mu)$$

is well-defined. Based on $v^*(c, \mu)$, I can study the limit payoffs as the prior probability of semi-rationality vanishes. Consider the correspondence $cl(Gr(v^*))$ defined by the closure of the graph of v^* , then

$$V^*(c) = cl(Gr(v^*))(c, \mathbf{0})$$

is the set of interest. To characterize $V^*(c)$, I first decompose the space of priors $[0, 1]^N$ into equivalence classes as follows. Let $\mu_0 = 0$. For each $\mu \in [0, 1]^N$, let $\omega(\mu) \in \mathfrak{R}^{(1+|N|) \times (1+|N|)}$ record the order of $\{\mu_0, \mu_1, \mu_2, \mu_3\}$ such that for all $k, k' \in \{0\} \cup N$,

$$\omega_{k,k'}(\mu) = (\mathbf{1}_{0 < \min\{\mu_k, \mu_{k'}\}} - \mathbf{1}_{0 = \min\{\mu_k, \mu_{k'}\}}) \cdot \text{sgn}(\mu_k - \mu_{k'}).$$

Then a pair of priors $\mu, \mu' \in [0, 1]^N$ is said to be ω -equivalent if $\omega(\mu) = \omega(\mu')$. Let $[\mu]_\omega$ denote the ω -equivalence class containing μ .

The second main result states that $V^*(c)$ is independent of c , efficient, and fully characterized by ω -equivalence classes.

Theorem 2 *For all $c \in (\frac{1}{2}, 1)$, $V^*(c) = V^* \subset X$. For each $x \in V^*$, there is a sequence $\mu^n \rightarrow \mathbf{0}$ such that for all n , $\mu^n \in [\mu^1]_\omega$, and $x = \lim_{n \rightarrow \infty} v^*(c, \mu^n)$. If two sequences $(\mu^{n,1})$ and $(\mu^{n,2})$ belong to the same ω -equivalence class, then $\lim_{n \rightarrow \infty} v^*(c, \mu^{n,1}) = \lim_{n \rightarrow \infty} v^*(c, \mu^{n,2})$. For all $i, j \in N$, $\text{sgn}(x_i - x_j) = -\omega_{i,j}(\mu^1)$.*

Intuitively, $V^*(c)$ is independent of c as c enters payoff when the rational type of an agent needs to assess situations with presence of semi-rational types. As the prior probability of semi-rationality vanishes, the importance of these situations also vanishes. Efficiency of V^* is not surprising given the Coasian feature of the bargaining process. The remaining part of Theorem 2 may be more interesting. If lower positive prior probability of semi-rationality is regarded as weaker reputation of commitment, then Theorem 2 implies that, a rational type who has weaker reputation obtains a larger share of the dollar: if μ is close to $\mathbf{0}$ and $0 < \mu_i < \mu_j$, then, $\text{sgn}(v_i^*(c, \mu) - v_j^*(c, \mu)) = 1$ and $v_i^*(c, \mu) > v_j^*(c, \mu)$. However, although weaker reputation compared to other agents delivers higher payoff, revealing rationality prior to the start of the bargaining process is not beneficial: if μ is close to $\mathbf{0}$ and $0 = \mu_i < \mu_j$, then $\text{sgn}(v_i^*(c, \mu) - v_j^*(c, \mu)) = -1$ and $v_i^*(c, \mu) < v_j^*(c, \mu)$. As the flipside of threat of exclusion, rational types are also tempted by inclusion. Note that a rational proposer faces a tradeoff between probability of agreement and terms of agreement. It turns out that, when comparing two agents with private information, probability of agreement is the dominating concern and he prefers to approach the agent with weaker reputation; but when comparing an agent with weak reputation and an revealed rational agent, terms of agreement is the dominating concern and he prefers to approach the agent with weak reputation as he can take advantage of threat of denouncement. Therefore, the agent with the weakest reputation can guarantee inclusion in all winning coalitions formed on equilibrium path. With the protection on terms of agreement by uniform recognition probabilities, he becomes the most advantageous bargainer. This is in sharp contrast to the bilateral model of Abreu and Gul (2000), in which a similar analysis on limit equilibrium payoffs shows that the rational type with stronger reputation obtains higher payoff. Another message of Theorem 2 is that the complete information stationary equilibrium payoff $(\frac{1}{3}, \frac{1}{3}, \frac{1}{3})$ is robust to symmetric perturbation by reputation: if $\mu^n \rightarrow \mathbf{0}$ is such that $\mu^n \in [\mu^1]_\omega$ and $\lim_{n \rightarrow \infty} v^*(c, \mu^n) = (\frac{1}{3}, \frac{1}{3}, \frac{1}{3})$, then for all n , $\mu_1^n = \mu_2^n = \mu_3^n$.

4 Equilibrium construction

In this section, I construct a symmetric stationary perfect Bayesian equilibrium from the ground of complete information bargaining to three-sided reputation bargaining by gradually increasing sides of private information. Then we prove our main theorems based on this construction.

Denote by $w_i^\delta(j, c, \mu)$ the post-recognition payoff for the rational type of agent i when agent j proposes, and let $W^\delta(c, \mu)$ be the $|N| \times |N|$ matrix with entries defined by $W_{ji}^\delta(c, \mu) = w_i^\delta(j, c, \mu)$. Denote by $v^\delta(c, \mu)$ the equilibrium payoff vector for rational types, and we have $v^\delta(c, \mu) = \frac{1}{3} \cdot \mathbf{1}^T W^\delta(c, \mu)$. I say that agent i actively concedes to agent j if agent i proposes an alternative x with $x_j \geq c$, and agent i passively concedes to agent j if agent i accepts an alternative in X_ϵ^j proposed by agent j . For each disjoint pair of $N', N'' \subseteq N$, let

$$\overline{\mathcal{M}}^{N'; N''} = \left\{ \mu \in [0, 1]^N \mid \mu_{N'} = \mathbf{0}_{N'}, \mu_{N''} = \mathbf{1}_{N''} \right\}$$

be the set of beliefs that agents in N' are rational with probability one, while agents in N'' are semi-rational with probability one. Note that $\overline{\mathcal{M}}^{N'; N''}$ is a face of $[0, 1]^N$. To facilitate presentation of belief evolution when agents take completely revealing actions, let

$proj_{N';N''}\mu$ be the projection of μ on $\overline{\mathcal{M}}^{N';N''}$. Also let

$$\mathcal{M}^{N';N''} = \left\{ \mu \in \overline{\mathcal{M}}^{N';N''} \mid \forall i \in N \setminus (N' \cup N''), \mu_i \in (0, 1) \right\}$$

be the set of beliefs that agents not in N' and N'' have private information.

4.1 Complete information

There are four cases of complete information bargaining depending on different numbers of revealed rational types. For all $j \in N$, I refer to the type with prior probability one of agent j simply as agent j . Since a relevant $\mathcal{M}^{N';N''}$ is a singleton, let $\mu^{N''}$ be its only element.

For the case with three revealed rational agents, Eraslan (2002) shows that the unique stationary subgame perfect equilibrium payoff is $(\frac{1}{3}, \frac{1}{3}, \frac{1}{3})$, and the payoff to a proposer is $1 - \frac{1}{3}\delta$. By symmetry, equilibrium strategies are also unique, i.e., a proposer mixes with uniform probabilities over offering $\frac{1}{3}\delta$ to one of the other agents and obtains immediate acceptances. Let $W^\delta(c, \mathbf{0}) = (1 - \frac{1}{3}\delta) \mathbf{I} + \frac{1}{6}\delta (\mathbf{1}^{|N| \times |N|} - \mathbf{I})$.

Proposition 3 (4.1.1) *Fix $c \in (\frac{1}{2}, 1)$ and $\mu = \mathbf{0}$. There is a unique symmetric stationary equilibrium in which bargaining ends immediately. Equilibrium payoffs for rational types are given by $v^\delta(c, \mathbf{0}) = \frac{1}{3} \cdot \mathbf{1}^T W^\delta(c, \mathbf{0})$*

For the case with two revealed rational types, WLOG let $\mu = \mu^{\{3\}}$. For each $i \in \{1, 2\}$, denote the revealed rational agent different from i by $-i$. First it is shown that neither agent 1 nor agent 2 actively concedes to agent 3.

Lemma 1 *Fix $c \in (\frac{1}{2}, 1)$ and $\mu = \mu^{\{3\}}$. For all $i \in \{1, 2\}$, $w_i^\delta(i, c, \mu^{\{3\}}) > 1 - c$.*

Proof. Since agent i as proposer can guarantee a payoff of $1 - c$ by conceding to agent 3, we have $w_i^\delta(i, c, \mu^{\{3\}}) \geq 1 - c$. Suppose toward a contradiction that $w_i^\delta(i, c, \mu^{\{3\}}) = 1 - c$. Then optimality of proposals for agent i implies that $v_{-i}^\delta(c, \mu^{\{3\}}) \geq \frac{c}{\delta}$, otherwise proposing to agent $-i$ is a profitable deviation. By symmetry, $v_i^\delta(c, \mu^{\{3\}}) \geq \frac{c}{\delta}$. Let $\hat{x}^i$ be such that $\hat{x}_i^i = 1$. Since $c > \frac{1}{2}$, we have $\delta v_i^\delta(c, \mu^{\{3\}}) \geq c > 1 - c$, which implies that agent i can obtain $\delta v_i^\delta(c, \mu^{\{3\}}) > 1 - c$ by proposing $\hat{x}^i$ which is rejected by agent $-i$ as $\delta v_{-i}^\delta(c, \mu^{\{3\}}) \geq c > \hat{x}_{-i}^i$ and by agent 3 as $\hat{x}_3^i < c$. But then $w_i^\delta(i, c, \mu^{\{3\}}) \geq \delta v_i^\delta(c, \mu^{\{3\}}) > 1 - c$, which is a contradiction. ■

Second it is shown that there is an immediate agreement between agent 1 and agent 2 whenever either of them proposes.

Lemma 2 *Fix $c \in (\frac{1}{2}, 1)$ and $\mu = \mu^{\{3\}}$. For all $i \in \{1, 2\}$, $w_i^\delta(i, c, \mu^{\{3\}}) = 1 - \delta v_{-i}^\delta(c, \mu^{\{3\}}) > \delta v_i^\delta(c, \mu^{\{3\}})$ and $w_i^\delta(-i, c, \mu^{\{3\}}) = \delta v_i^\delta(c, \mu^{\{3\}})$.*

Proof. Suppose toward a contradiction that $1 - \delta v_{-i}^\delta(c, \mu^{\{3\}}) \leq \delta v_i^\delta(c, \mu^{\{3\}})$, then by symmetry, $\delta v_i^\delta(c, \mu^{\{3\}}) \geq \frac{1}{2} > 1 - c$. Since neither agent 1 nor agent 2 actively concede to agent 3, for all $j \in N$, $w_i^\delta(j, c, \mu^{\{3\}}) = \delta v_i^\delta(c, \mu^{\{3\}})$. But then $v_i^\delta(c, \mu^{\{3\}}) = \delta v_i^\delta(c, \mu^{\{3\}}) = 0$, which is a contradiction. Therefore $1 - \delta v_{-i}^\delta(c, \mu^{\{3\}}) > \delta v_i^\delta(c, \mu^{\{3\}})$, and by symmetry $\delta v_i^\delta(c, \mu^{\{3\}}) < \frac{1}{2} < c$, verifying that neither agent 1 nor agent 2 actively concede to agent

3. Since agent i as proposer can guarantee a payoff of $1 - \delta v_{-i}^\delta(c, \mu^{\{3\}}) - \varepsilon$ for all $\varepsilon > 0$, it follows that $w_i^\delta(i, c, \mu^{\{3\}}) = 1 - \delta v_{-i}^\delta(c, \mu^{\{3\}})$. Note that $w_i^\delta(i, c, \mu^{\{3\}}) = 1 - \delta v_{-i}^\delta(c, \mu^{\{3\}}) > \delta v_i^\delta(c, \mu^{\{3\}})$ implies immediate agreement. ■

Given these two lemmas, it can be shown that when δ is close to 1, there is an equilibrium in which neither of agent 1 and agent 2 passively concede to agent 3, and thus agent 3 is essentially excluded from bargaining. In this equilibrium, the best response of the concealed rational type of agent 3 is to reveal rationality when he proposes. Let $W^\delta(c, \mu^{\{3\}})$ be defined as follows: (1) for all $i \in \{1, 2\}$, let

$$\begin{aligned} w_i^\delta(i, c, \mu^{\{3\}}) &= 1 - \delta v_{-i}^\delta(c, \mu^{\{3\}}) \\ w_i^\delta(-i, c, \mu^{\{3\}}) &= \delta v_i^\delta(c, \mu^{\{3\}}) \\ w_i^\delta(3, c, \mu^{\{3\}}) &= \delta v_i^\delta(c, \mu^{\{3\}}), \end{aligned}$$

with the first two equalities representing the equilibrium plays that agent 1 and agent 2 seek agreement with each other when either of them proposes, and the last equality representing that a proposal by agent 3 is rejected with probability one; (2) for the concealed rational type of agent 3, let

$$\begin{aligned} w_3^\delta(i, c, \mu^{\{3\}}) &= 0 \\ w_3^\delta(3, c, \mu^{\{3\}}) &= w_3^\delta(3, c, \mathbf{0}) \end{aligned}$$

with the first equality representing that his exclusion when either of agent 1 and agent 2 proposes, and the last equality representing his active revelation of rationality.

Proposition 4 (4.1.2) *Fix $c \in (\frac{1}{2}, 1)$ and $\mu = \mu^{\{3\}}$. There is $\delta^*(c, \mu^{\{3\}}) \in (0, 1)$ such that for all $\delta > \delta^*(c, \mu^{\{3\}})$, there is a symmetric stationary perfect Bayesian equilibrium. In this equilibrium, bargaining ends immediately if either of agent 1 and agent 2 is recognized as the proposer, and there is no concession to agent 3. The concealed rational type of agent 3 would actively reveal rationality. Equilibrium payoffs for rational types are given by $v^\delta(c, \mu^{\{3\}}) = \frac{1}{3} \cdot \mathbf{1}^T W^\delta(c, \mu^{\{3\}})$.*

Proof. Suppose there is no passive concession, then for all $i \in \{1, 2\}$, $w_i^\delta(3, c, \mu^{\{3\}}) = \delta v_i^\delta(c, \mu^{\{3\}})$. By Lemma 2, we have $v_i^\delta(c, \mu^{\{3\}}) = \frac{1}{3-\delta}$. Then rejection of $x \in X_c^3$ by agent i is justified if $\delta v_i^\delta(c, \mu^{\{3\}}) = \frac{\delta}{3-\delta} \geq 1-c$, or equivalently, $\delta > \frac{3(1-c)}{2-c}$. Note that $\frac{3(1-c)}{2-c} \in (0, 1)$ is equivalent to $c \in (\frac{1}{2}, 1)$. As for the concealed rational agent 3, if he maintains concealed, then by stationarity $v_3^\delta(c, \mu^{\{3\}}) = \frac{1}{3}\delta v_3^\delta(c, \mu^{\{3\}}) = 0$; but if he actively reveals rationality, he obtains $w_3^\delta(3, c, \mathbf{0}) = 1 - \frac{1}{3}\delta > 0$. Active revelation of rationality is feasible if $1 - \frac{1}{3}\delta \neq c$. For $c \leq \frac{2}{3}$, $1 - \frac{1}{3}\delta > c$ for all $\delta \in (0, 1)$; while for $c > \frac{2}{3}$, $1 - \frac{1}{3}\delta < c$ for all $\delta > 3(1-c)$. Therefore, we can set $\delta^*(c, \mu^{\{3\}}) = \frac{3(1-c)}{2-c}$ for $c \leq \frac{2}{3}$ and set $\delta^*(c, \mu^{\{3\}}) = 3(1-c)$ for $c > \frac{2}{3}$. Equilibrium payoffs for rational types are then straightforward to compute. ■

Although there may be other equilibria, the one selected in Proposition 4.1.2 maximizes not only exclusion of the semi-rational agent but also the payoffs for the revealed rational types. To see this, if there is passive concession in equilibrium, then the equilibrium payoff for a revealed rational type is at most $\frac{1-c}{\delta}$ which is lower than $v_1^\delta(c, \mu^{\{3\}})$ when δ is close to 1. In certain sense, the revealed rational types face a coordination problem in excluding the

semi-rational agent. In the equilibrium selected in Proposition 4.1.2, the revealed rational types successfully coordinate to achieve the efficient outcome for them, which is one reason why I am interested in it.

For the case with only one revealed rational type, WLOG let $\mu = \mu^{\{2,3\}}$. For each $i \in \{2,3\}$, denote the semi-rational agent different from i by $-i$. Based on Proposition 4.1.2, there is an equilibrium in which bargaining ends immediately and concealed rational types maintain concealed. Immediate agreement entails that agent 1 concedes immediately. Let $W^\delta(c, \mu^{\{2,3\}})$ be defined as follows: (1) for all $j \in N$, let

$$w_1^\delta(j, c, \mu^{\{2,3\}}) = 1 - c$$

which represents concession by agent 1; (2) for all $i \in \{2,3\}$, let

$$\begin{aligned} w_i^\delta(1, c, \mu^{\{2,3\}}) &= \frac{1}{2}c \\ w_i^\delta(i, c, \mu^{\{2,3\}}) &= c \\ w_i^\delta(-i, c, \mu^{\{2,3\}}) &= 0 \end{aligned}$$

with the first equality representing the equilibrium plays that agent 1 mixes with uniform probabilities over concession to one of the other agents, and the last two equalities representing passive concession of agent 1 and that concealed rational types maintain concealed as proposers.

Proposition 5 (4.1.3) *Fix $c \in (\frac{1}{2}, 1)$ and $\mu^{\{2,3\}}$. There is $\delta^*(c, \mu^{\{2,3\}}) \in (0, 1)$ such that for all $\delta > \delta^*(c, \mu^{\{2,3\}})$, there is a unique symmetric stationary perfect Bayesian equilibrium. In this equilibrium, bargaining ends immediately in concessions of agent 1 to other agents. Equilibrium payoffs for rational types are given by $v^\delta(c, \mu^{\{2,3\}}) = \frac{1}{3} \cdot \mathbf{1}^T W^\delta(c, \mu^{\{2,3\}})$*

Proof. First we show that agent 1 passively concedes. Suppose not, then $\delta v_1^\delta(c, \mu^{\{2,3\}}) \geq 1 - c$ and thus $v_1^\delta(c, \mu^{\{2,3\}}) = \delta v_1^\delta(c, \mu^{\{2,3\}}) = 0$, which is a contradiction. Second we show that agent 1 actively concedes. Suppose not, then $w_1^\delta(1, c, \mu^{\{2,3\}}) = \delta v_1^\delta(c, \mu^{\{2,3\}}) \geq 1 - c$ and

$$\begin{aligned} v_1^\delta(c, \mu^{\{2,3\}}) &= \frac{1}{3} \delta v_1^\delta(c, \mu^{\{2,3\}}) + \frac{2}{3} (1 - c) \\ &= \frac{2}{3 - \delta} (1 - c) \\ &< \frac{1 - c}{\delta}, \end{aligned}$$

which is a contradiction. Since agent 1 passively concedes, when a concealed rational type proposes, he obtains c by maintaining concealed, and he obtains $1 - \frac{\delta}{3 - \delta}$ by revealing rationality. When $\delta > \frac{3(1-c)}{2-c}$, $c > 1 - \frac{\delta}{3 - \delta}$ and a concealed rational type as proposer maintains concealed. When a concealed rational type is conceded to, he immediately accepts and obtains c , as rejection would reveal rationality and gives $\frac{\delta}{3 - \delta} < c$. By symmetry, agent 1 mixes with uniform probabilities over offering c to one of the other agents. Then the equilibrium payoffs follow from simple calculation. ■

Note that the equilibrium outcome is unique in Proposition 4.1.3, but the equilibrium payoffs for rational types depends on the equilibrium selected in Proposition 4.1.2. If an equilibrium with passive concession is selected at $\mu^{\{2\}}$ and $\mu^{\{3\}}$, then at $\mu^{\{2,3\}}$, the concealed rational types may prefer to actively reveal rationality.

For the case with no revealed rational types, i.e., $\mu = \mathbf{1}$, bargaining never ends in agreement and for all $i \in N$, the semi-rational type of i is "indifferent" to all alternatives in X_c^i . Although the specification of proposing behavior of semi-rational types does not affect bargaining outcome, it does affects the payoffs for concealed rational types in a subtle way. For example, if each semi-rational agent mixes with uniform probabilities over $(c; 1 - c, 0)$ and $(c; 0, 1 - c)$, the payoff for the concealed rational type of agent i would be given by

$$\begin{aligned} v_i(\delta) &= \frac{1}{3}(1 - c) + \frac{2}{3} \left[\frac{1}{2}(1 - c) + \frac{1}{2}\delta v_i(\delta) \right] \\ &= \frac{2}{3 - \delta}(1 - c); \end{aligned}$$

however, if each semi-rational agent proposes $(c; \frac{1}{2}(1 - c), \frac{1}{2}(1 - c))$ with probability one, the payoff for the concealed rational type of agent i would be

$$\begin{aligned} v'_i(\delta) &= \frac{1}{3}(1 - c) + \frac{2}{3}\delta v'_i(\delta) \\ &= \frac{1}{3 - 2\delta}(1 - c). \end{aligned}$$

Note that $v_i(1) = v'_i(1)$, but for all $\delta \in (0, 1)$, $v_i(\delta) > v'_i(\delta)$. This vanishing difference turns out to be important when analyzing cases with two- or three-sided private information. To maximize threat of denouncement, the latter specification is selected. And thus all $i, j \in N$

$$w_i^\delta(j, c, \mathbf{1}) = \mathbf{1}_{i=j} \cdot (1 - c) + \mathbf{1}_{i \neq j} \cdot \delta v_i^\delta(c, \mathbf{1}).$$

4.2 One-sided private information

There are three cases with one-sided private information depending on different numbers of revealed rational types.

For the case with two revealed rational types, WLOG let $\mu \in \mathcal{M}^{\{1,2\};\emptyset}$. In this case, one can exploit the threat of denouncement given by Proposition 4.1.2, and shows that when δ is close to 1, there is a separating equilibrium. For each $i \in \{1, 2\}$, denote the revealed rational type different from i by $-i$. When agent i proposes, he screens agent 3 with the threat to enter the continuation with posterior $\mu^{\{3\}}$, and the rational type of agent 3 accepts immediately. The proposal stage of agent 3 is also separating. The rational type of agent 3 reveals rationality and the game proceeds with posterior $\mathbf{0}$, while the semi-rational type proposes alternatives in X_c^3 and the game proceeds with posterior $\mu^{\{3\}}$. Then let $W^\delta(c, \mu)$ be defined as follows: for all $i \in \{1, 2\}$,

$$\begin{aligned} w_i^\delta(i, c, \mu) &= (1 - \mu_3) \left(1 - \delta v_3^\delta(c, \mu^{\{3\}}) \right) + \mu_3 \delta v_i^\delta(c, \mu^{\{3\}}) \\ w_i^\delta(-i, c, \mu) &= (1 - \mu_3) \cdot 0 + \mu_3 \delta v_i^\delta(c, \mu^{\{3\}}) \\ w_i^\delta(3, c, \mu) &= (1 - \mu_3) \cdot \frac{1}{2} \delta v_i^\delta(c, \mathbf{0}) + \mu_3 \delta v_i^\delta(c, \mu^{\{3\}}) \end{aligned}$$

and for the rational type of agent 3, let

$$\begin{aligned} w_3^\delta(i, c, \mu) &= \delta v_3^\delta(c, \mu^{\{3\}}) \\ w_3^\delta(3, c, \mu) &= w_3^\delta(3, c, \mathbf{0}). \end{aligned}$$

The previous discussion is summarized into the following proposition.

Proposition 6 (4.2.1) *Fix $c \in (\frac{1}{2}, 1)$ and $\mu \in \mathcal{M}^{\{1,2\};\emptyset}$. There is $\delta^*(c, \mu) \in (0, 1)$ such that for all $\delta > \delta^*(c, \mu)$, there is a separating symmetric stationary perfect Bayesian equilibrium. Equilibrium payoffs for rational types are given by $v^\delta(c, \mu) = \frac{1}{3} \cdot \mathbf{1}^T W^\delta(c, \mu)$.*

The idea in this construction is to maximize the informativeness of actions taken by agents with private information and minimize the number of rounds of screening, and it has very similar counterparts in proving most of the propositions on equilibrium construction below. Here I provide a sketch of proof, and the details are in Appendix. (1) Consider the proposal stage of agent 3. If agent 3 makes a proposal in X_c^3 , the belief is updated to $\mu^{\{3\}}$, then the rational agent 3 would prefer to reveal rationality and obtain $w_3^\delta(3, c, \mathbf{0})$ by Proposition 4.1.2. (2) Consider the proposal stage of agent 1. Since $\delta v_3^\delta(c, \mu^{\{3\}}) < \frac{1}{3} < c$, agent 1 can propose x with $x_3 \in [\delta v_3^\delta(c, \mu^{\{3\}}), c)$, and if agent 3 rejects, the belief is updated to $\mu^{\{3\}}$. Then it is plausible that the rational agent 3 would accept x given the continuation plays associated with $\mu^{\{3\}}$. It remains to verify that the optimal proposal that agent 1 would make is $(1 - \delta v_3^\delta(c, \mu^{\{3\}}), 0, \delta v_3^\delta(c, \mu^{\{3\}}))$ and it can be decomposed into four steps: active concession to agent 3 is dominated; seeking agreement with agent 2 is dominated; non-serious proposals are dominated; $v_3^\delta(c, \mu^{\{3\}}) < v_3^\delta(c, \mu)$ and thus all proposals x with $x_3 < \delta v_3^\delta(c, \mu^{\{3\}})$ are either seeking agreement with agent 2 or non-serious. The last statement implies that there is no need for multiple rounds of screening and thus it is much easier to write down explicitly the equilibrium payoffs. Starting from this proposition, verification of some inequalities for a fixed δ turns out to be much more complicated than to examine the limits as $\delta \rightarrow 1$. If an inequality is strict in the limit, then we can claim that it holds for δ close 1. It is in this way that patient rational types improve tractability.

For the case with one revealed rational type, WLOG let $\mu \in \mathcal{M}^{\{1\};\{3\}}$. In this case, one can exploit the threat of denouncement given by Proposition 4.1.3, and shows that when δ is close to 1, there is an almost-separating equilibrium. Almost-separation means that when bargaining ends, the realized posterior is close to complete information. A fully separating equilibrium is impossible as the proposal stage of agent 2 cannot be separating with the continuation given Proposition 4.1.3. To see this, suppose toward a contradiction that the proposal stage of agent 2 is separating. When δ is close to 1, the rational type of agent 2 obtains approximately $\frac{1}{2}$ by revealing rationality, but if he proposes $(1 - c, c, 0)$, he is regarded as semi-rational and obtains $c > \frac{1}{2}$ as agent 1 immediately concedes, which is a contradiction. Another complication in this case is that the proposals by agent 3 can also be screening when the claim is moderate. Then to balance agents' incentives, we let the probability that agent 3 seeks agreement with agent 1 to be endogenous.

Here I present the key elements of equilibrium construction. The verification of the construction and other details are in Appendix. (1) First, let $\hat{m}^\delta$ be the solution to

$$\begin{aligned} \frac{1}{\delta}(1 - c) &= \frac{1}{3} \left[(1 - \hat{m}^\delta) \left(1 - \delta v_2^\delta(c, \mu^{\{2,3\}}) \right) + \hat{m}^\delta \delta v_1^\delta(c, \mu^{\{2,3\}}) \right] \\ &\quad + \frac{2}{3}(1 - c). \end{aligned}$$

The idea is to make agent 1 indifferent between acceptance and rejection of $(1 - c, c, 0)$, when he believes that agent 2 is semi-rational with probability $\hat{m}^\delta$ and he plans to screen if he is recognized in the next period but to passively concede if other agents are recognized. Then we can exploit mixed voting to control the payoff for the rational type of agent 2 when he mimicks the semi-rational type. It is easy to see that $\lim_{\delta \rightarrow 1} \hat{m}^\delta = 1$. Given the strategy of agent 1, neither of the other rational types would actively reveal rationality at $(0, \hat{m}^\delta, 1)$ when δ is close to 1. Accordingly, we define $W^\delta(c, (0, \hat{m}^\delta, 1))$ and $v^\delta(c, (0, \hat{m}^\delta, 1))$. (2) Next consider δ close to 1 such that $\mu_2 < \hat{m}^\delta$. Let $w(1, c, \mu)$ be associated with the plays that agent 1 screens agent 2 by the threat to enter the continuation with posterior $\mu^{\{2,3\}}$. Since the proposal stage of agent 2 is semi-pooling, the rational type of agent 2 is indifferent between revealing rationality and mimicking the semi-rational type. Let r^δ be the solution to

$$1 - \delta v_1^\delta(c, \mu^{\{3\}}) = (1 - r^\delta)c + r^\delta \delta v_2^\delta(c, (0, \hat{m}^\delta, 1)).$$

and

$$\begin{aligned} w_1^\delta(2, c, \mu) &= \left(1 - \frac{\mu_2}{\hat{m}^\delta}\right) \delta v_1^\delta(c, \mu^{\{3\}}) + \frac{\mu_2}{\hat{m}^\delta} (1 - c) \\ w_2^\delta(2, c, \mu) &= 1 - \delta v_1^\delta(c, \mu^{\{3\}}) \\ w_3^\delta(2, c, \mu) &= \frac{\mu_2}{\hat{m}^\delta} r^\delta \delta v_3^\delta(c, (0, \hat{m}^\delta, 1)). \end{aligned}$$

Note that $\frac{\mu_2}{\hat{m}^\delta}$ is the total probability that agent 2 behaves as semi-rational if by observing such an action, the belief updates from μ_2 to $\hat{m}^\delta$ according to Bayes' rule. If $(0, 1 - c, c)$ is a screening proposal, i.e., $1 - c > \delta v_2^\delta(c, \mu^{\{2,3\}})$, the rational type of agent 2 will passively concede to agent 3. Consider the system

$$\begin{aligned} w_1^\delta(3, c, \mu) &= p^\delta \max\{\delta v_1^\delta(c, \mu), 1 - c\} + (1 - p^\delta) \mu_2 \delta v_1^\delta(c, \mu^{\{2,3\}}) \\ v_1^\delta(c, \mu) &= \sum_{j \in N} \frac{1}{3} w_1(j, c, \mu) \\ p^\delta &= 0, \text{ if } \delta v_1^\delta(c, \mu) > 1 - c. \end{aligned}$$

p^δ is the probability that agent 3 seeks agreement with agent 1 and it is uniquely determined by the system above. Given p^δ , let

$$w_2^\delta(3, c, \mu) = p^\delta \mu_2 \delta v_2^\delta(c, \mu) + (1 - p^\delta) (1 - c),$$

which says that if agent 1 passively concedes, he concedes with probability μ_2 to balance the incentive for agent 3. If $(0, 1 - c, c)$ is not a screening proposal, i.e., $1 - c \leq \delta v_2^\delta(c, \mu^{\{2,3\}})$, the equilibrium plays are simpler, neither of agent 1 and agent 2 concedes to agent 3, and we have for all $i \in \{1, 2\}$,

$$w_i^\delta(3, c, \mu) = \delta v_i^\delta(c, \mu).$$

By Proposition 4.2.1, it can be shown that the concealed rational type of agent 3 will actively reveal rationality at μ and thus

$$w_3^\delta(3, c, \mu) = w_3^\delta(3, c, proj_{1,3;\emptyset}\mu).$$

Thus $W^\delta(c, \mu)$ is defined, and the previous discussion is summarized into the following proposition.

Proposition 7 (4.2.2) Fix $c \in (\frac{1}{2}, 1)$ and $\mu \in \mathcal{M}^{\{1\};\{3\}}$. There is $\delta^*(c, \mu) \in (0, 1)$ such that for all $\delta > \delta^*(c, \mu)$, there is an almost-separating symmetric stationary perfect Bayesian equilibrium. In this equilibrium, the concealed rational type of agent 3 actively reveals rationality if he is recognized first. Equilibrium payoffs for rational types are given by $v^\delta(c, \mu) = \frac{1}{3} \cdot \mathbf{1}^T W^\delta(c, \mu)$.

For the case with no revealed rational types, WLOG let $\mu \in \mathcal{M}^{\emptyset;\{1,2\}}$. For all $i \in \{1, 2\}$, let $-i$ be the other semi-rational agent. It is easy to see that the rational type of agent 3 concedes immediately, and the concealed rational types would seek agreements with the rational type of agent 3 if they are recognized first. If the concealed rational type of agent 1 actively reveal rationality, the game proceeds with posterior $proj_{1;2}\mu$ and agents play as described in Proposition 4.2.2. If the concealed rational type of agent 1 maintains concealed, he obtains c if agent 3 is rational, but the game would enter the continuation with posterior $\mathbf{1}$ if agent 3 is semi-rational. Thus whether a concealed rational agent would actively reveal rationality depends on model parameters. $W^\delta(c, \mu)$ can be defined as follows: for all $i \in \{1, 2\}$, let

$$w_3^\delta(i, c, \mu) = w_3^\delta(3, c, \mu) = 1 - c$$

and

$$\begin{aligned} w_i^\delta(i, c, \mu) &= \max \{w_i^\delta(i, c, proj_{i;-i}\mu), (1 - \mu_3)c + \mu_3\delta v_i^\delta(c, \mathbf{1})\} \\ w_i^\delta(-i, c, \mu) &= \mu_3\delta v_i^\delta(c, \mathbf{1}) \\ w_i^\delta(3, c, \mu) &= (1 - \mu_3) \cdot \frac{1}{2}c + \mu_3\delta v_i^\delta(c, \mathbf{1}). \end{aligned}$$

Proposition 8 (4.2.3) Fix $c \in (\frac{1}{2}, 1)$ and $\mu \in \mathcal{M}^{\emptyset;\{1,2\}}$. There is $\delta^*(c, \mu) \in (0, 1)$ such that for all $\delta > \delta^*(c, \mu)$, there is a separating symmetric stationary perfect Bayesian equilibrium. In this equilibrium, agent 3 concedes immediately; a concealed rational type of agent $i \in \{1, 2\}$ actively reveals rationality if $w_i^\delta(i, c, proj_{i;-i}\mu) > (1 - \mu_3)c + \mu_3\delta v_i^\delta(c, \mathbf{1})$. Equilibrium payoffs for rational types are given by $v^\delta(c, \mu) = \frac{1}{3} \cdot \mathbf{1}^T W^\delta(c, \mu)$.

4.3 Two-sided private information

There are two cases with two-sided private information depending on whether the agent without private information is rational or not. WLOG, let agent 1 be the agent without private information. And for each $i \in \{2, 3\}$, let $-i$ be the other agent with private information.

For the case in which agent 1 is revealed rational, fix $\mu \in \mathcal{M}^{\{1\};\emptyset}$. Based on Proposition 4.2.1 and Proposition 4.2.2, it can be shown that when δ is close to 1, there is an almost-separating equilibrium. In this equilibrium, if agent 1 is recognized first, he screens the agent i with weaker reputation by the threat to enter the continuation with posterior $proj_{1;i}\mu$. By Proposition 4.2.2, if agent $i \in \{2, 3\}$ is recognized first, the rational type reveals rationality and the game proceeds with posterior $proj_{1;i;\emptyset}\mu$, while the semi-rational type behaves as described in Proposition 4.2.2 and the game proceeds with posterior $proj_{1;i}\mu$. For agent 1, the payoff from screening agent i is

$$(1 - \mu_i) (1 - \delta v_i^\delta(c, proj_{1;i}\mu)) + \mu_i \delta v_1^\delta(c, proj_{1;i}\mu).$$

By Proposition 4.2.2, $1 - \delta v_i^\delta(c, \text{proj}_{1;i}\mu) < 1 - \delta v_i^\delta(c, \text{proj}_{1;-i}\mu)$ if and only if $\mu_i < \mu_{-i}$, i.e., the probability of immediate agreement is higher with the agent with weaker reputation but the term of agreement is worse. It turns out that the concern about probability of agreement dominates. When $\mu_i = \mu_{-i}$, symmetry requires that agent 1 mixes with uniform probabilities over seeking agreements with agent 2 or agent 3. Then it is straightforward to define $W^\delta(c, \mu)$. The details of $W^\delta(c, \mu)$ and the verification of the construction are in Appendix.

Proposition 9 (4.3.1) *Fix $c \in (\frac{1}{2}, 1)$ and $\mu \in \mathcal{M}^{\{1\};\emptyset}$. There is $\delta^*(c, \mu) \in (0, 1)$ such that for all $\delta > \delta^*(c, \mu)$, there is an almost-separating symmetric stationary perfect Bayesian equilibrium. In this equilibrium, if agent 1 is recognized first, he screens the agent with weaker reputation; for all $i \in \{2, 3\}$, if agent i is recognized, the rational type of agent i reveals rationality and screens agent $-i$. Equilibrium payoffs for rational types are given by $v^\delta(c, \mu) = \frac{1}{3} \cdot \mathbf{1}^T W^\delta(c, \mu)$.*

For the case in which agent 1 is semi-rational, fix $\mu \in \mathcal{M}^{\emptyset;\{1\}}$. There are two subcases depending on whether it is plausible that for all $i \in \{2, 3\}$, the proposal stage of agent i is separating. When the proposal stage of agent i is separating, the rational type of agent i is supposed to actively reveal rationality and the game proceeds with posterior $\text{proj}_{i;1}\mu$. Then he obtains a payoff of $w_i^\delta(i, c, \text{proj}_{i;1}\mu)$. However, if he mimicks the semi-rational type and forces concession by the rational type of agent $-i$, he obtains $(1 - \mu_{-i})c + \mu_{-i}\delta v_i^\delta(c, \mathbf{1})$. Thus if $w_i^\delta(i, c, \text{proj}_{i;1}\mu) > (1 - \mu_{-i})c + \mu_{-i}\delta v_i^\delta(c, \mathbf{1})$, a separating proposal stage of agent i is plausible. Given that the rational types would seek agreement with each other, they don't concede to agent 1. Then it is easy to show that the concealed rational type of agent 1 would actively reveal rationality. Let $W^\delta(c, \mu)$ be defined as follows: for all $i \in \{2, 3\}$, let

$$\begin{aligned} w_i^\delta(i, c, \mu) &= w_i^\delta(i, c, \text{proj}_{i;1}\mu) \\ w_i^\delta(-i, c, \mu) &= \delta v_i^\delta(c, \text{proj}_{i;1,-i}\mu) \\ w_i^\delta(3, c, \mu) &= \delta v_i^\delta(c, \mu) \end{aligned}$$

and

$$\begin{aligned} w_1^\delta(1, c, \mu) &= w_1^\delta(1, c, \text{proj}_{1;\emptyset}\mu) \\ w_1^\delta(i, c, \mu) &= \mu_{-i}\delta v_1^\delta(c, \text{proj}_{i;1,-i}\mu). \end{aligned}$$

If $w_i(i, c, \text{proj}_{i;1}\mu) < (1 - \mu_{-i})c + \mu_{-i}\delta v_i^\delta(c, \mathbf{1})$, i.e., a separating proposal stage of agent i is implausible, the construction is more complicated. An argument similar to Proposition 4.2.2 shows that the proposal stage of agent i is almost-separating. An additional layer of complication comes from the fact that the voting decision of agent $-i$ should also be semi-pooling. I postulate equilibrium strategies as follows. Suppose agent 1 is recognized first, neither agent 2 nor agent 3 concede. Suppose agent 2 is recognized first. The semi-rational type of agent 2 forces concession by agent 3 and proposes $(0, c, 1 - c)$. The rational type of agent 2 mixes between revealing rationality and mimicking the semi-rational type. The latter action induces posterior $(1, \hat{m}^{\delta,(1)}, \mu_3)$. At $(1, \hat{m}^{\delta,(1)}, \mu_3)$, the rational type of agent 3 is indifferent between acceptance and rejection. And by rejection, posterior is updated to $(1, \hat{m}^{\delta,(1)}, \hat{m}^{\delta,(2)})$. At $(1, \hat{m}^{\delta,(1)}, \hat{m}^{\delta,(2)})$, the rational type of agent 3 passively concedes to agent 1 and agent 2. Then accordingly, agent 1 and agent 2 would force concession

by agent 3.³ Now consider the proposal stage of agent 3 at $(1, \hat{m}^{\delta,(1)}, \hat{m}^{\delta,(2)})$ and suppose it is semi-pooling.⁴ By proposing $(0, 1 - c, c)$, posterior is updated to $(1, \hat{m}^{\delta,(1)}, \bar{m}^{\delta,(2)})$. At $(1, \hat{m}^{\delta,(1)}, \bar{m}^{\delta,(2)})$, the rational type of agent 2 is indifferent between acceptance and rejection. And by rejection, posterior is updated to $(1, \bar{m}^{\delta,(1)}, \bar{m}^{\delta,(2)})$. At $(1, \bar{m}^{\delta,(1)}, \bar{m}^{\delta,(2)})$, agent 2 and agent 3 force concession by one another; and agent 1 forces concession by agent 3. It can be verified that the strategies defined above is an equilibrium. Also both $\hat{m}^{\delta,(1)}$ and $\bar{m}^{\delta,(2)}$ converge to 1 as $\delta \rightarrow 1$, i.e., the first proposal stage of agent i is almost-separating. Most importantly, by Proposition 4.3.1, the concealed rational type of agent 1 would actively reveal rationality if he is recognized first, which implies that in the three-sided private information environment, the first proposal stage can be separating. Let $W^\delta(c, \mu)$ be defined according to the strategies postulated above, and details are in Appendix.

Proposition 10 (4.3.2) *Fix $c \in (\frac{1}{2}, 1)$ and $\mu \in \mathcal{M}^{\emptyset;\{1\}}$. There is $\delta^*(c, \mu) \in (0, 1)$ such that for all $\delta > \delta^*(c, \mu)$, there is a symmetric stationary perfect Bayesian equilibrium. In this equilibrium, for all $i \in \{2, 3\}$, the first proposal stage of agent i is almost-separating; and the concealed rational type of agent 1 would actively reveal rationality if he is recognized first. Equilibrium payoffs for rational types are given by $v^\delta(c, \mu) = \frac{1}{3} \cdot \mathbf{1}^T W^\delta(c, \mu)$.*

4.4 Three-sided private information

Finally, we are ready for the analysis of the case with three-sided private information, which is straightforward given Proposition 4.3.1 and Proposition 4.3.2. As mentioned above, a separating first proposal stage is plausible. If agent i is recognized first, the rational type of agent i would reveal rationality, and the game proceeds with posterior $proj_{i;\emptyset}\mu$. The proposals by the semi-rational type of agent i are rejected with probability one, and the game proceeds with posterior $proj_{\emptyset;i}\mu$. Then let $W^\delta(c, \mu)$ be defined as follows: for all $i, j \in N$ such that $i \neq j$,

$$\begin{aligned} w_i^\delta(i, c, \mu) &= w_i^\delta(i, c, proj_{i;\emptyset}\mu) \\ w_i^\delta(j, c, \mu) &= (1 - \mu_j) w_i^\delta(j, c, proj_{j;\emptyset}\mu) + \mu_j w_i^\delta(j, c, proj_{\emptyset;j}\mu). \end{aligned}$$

Proposition 11 (4.4) *Fix $c \in (\frac{1}{2}, 1)$ and $\mu \in \mathcal{M}^{\emptyset;\emptyset}$. There is $\delta^*(c, \mu) \in (0, 1)$ such that for all $\delta > \delta^*(c, \mu)$, there is a symmetric stationary perfect Bayesian equilibrium. In this equilibrium, the first proposal stage is separating. Equilibrium payoffs for rational types are given by $v^\delta(c, \mu) = \frac{1}{3} \cdot \mathbf{1}^T W^\delta(c, \mu)$.*

4.5 Proofs of main theorems

The two main theorems follow almost immediately from examining the equilibrium constructed by previous propositions.

³Note that the voting stage of agent 3 is now separating, which is plausible only if $1 - c \geq \delta v_3^\delta(c, (1, \hat{m}^{\delta,(1)}, 1))$. This is the point at which the specification of proposing behavior at $\mu = \mathbf{1}$ is important. It turns out that with the specification we adopted, separating voting of agent 3 at $(1, \hat{m}^{\delta,(1)}, \hat{m}^{\delta,(2)})$ is plausible but may not be so with other specifications.

⁴Some parameters of interest permit a separating proposal stage of agent 3 at $(1, \hat{m}^{\delta,(1)}, \hat{m}^{\delta,(2)})$, while others do not. However, a semi-pooling proposal stage of agent 3 is always plausible. To avoid further division into sub-subcases, we choose the latter approach.

4.5.1 Proof of Theorem 1

The existence part of Theorem 1 is obvious. To see that the Coasian part of Theorem 1 is true, we examine each case of priors. Recall that bargaining starts at period 0 and $T \geq 4$.

(1) Consider the cases with complete information. When there are odd number of revealed rational agents (see Proposition 4.1.1 and Proposition 4.1.3), bargaining ends immediately. When there are two revealed rational agents (see Proposition 4.1.2), bargaining ends if either of them is recognized, and thus the probability that bargaining ends by period T is $1 - (\frac{1}{3})^{T+1}$. When there are no revealed rational agents, bargaining never reaches agreement, which is a trivial case.

(2) Consider the cases with one-sided private information. When there are two revealed rational types (see Proposition 4.2.1), if the agent with private information is rational, then bargaining ends immediately; otherwise, the game enters a continuation described by Proposition 4.1.2. And thus bargaining ends by period T with probability at least $1 - (\frac{1}{3})^T$. When there is only one revealed rational types (see Proposition 4.2.2), if the revealed rational agent is recognized first and the agent with private information is rational, then bargaining ends immediately; otherwise, bargaining ends in the next period. If the agent with private information is recognized first, then bargaining ends immediately if he reveals rationality; otherwise, bargaining is guaranteed to end in the next two periods. It implies that bargaining ends by period T if either of these two agents is recognized during $0, \dots, T-2$, and thus the probability that bargaining ends by period T is at least $1 - (\frac{1}{3})^{T-1}$. When there is no revealed rational types (see Proposition 4.2.3), bargaining ends immediately if the agent with private information is rational; otherwise, bargaining never reaches agreement. And thus the probability that bargaining ends by period T is $1 - \prod_{i \in N} \mu_i$.

(3) Consider the cases with two-sided private information. When there is one revealed rational type (see Proposition 4.3.1), if there is no agreement in the first period, then the game enters either a continuation described by Proposition 4.1.2 or one described by Proposition 4.2.2. And thus the probability that bargaining ends by period T is at least $1 - (\frac{1}{3})^{T-2}$. When there is no revealed rational type (see Proposition 4.3.2), if an agent with private information is recognized and reveals rationality, then bargaining ends immediately or in the next period. If he behaves as semi-rational and at least one agent is rational, bargaining is guaranteed to end in the next three periods. Therefore, bargaining ends by period T , if at least one agent with private information is recognized during $0, \dots, T-3$ and at least one agent is rational, which implies that the probability that bargaining ends by period T is at least $(1 - \prod_{i \in N} \mu_i) [1 - (\frac{1}{3})]^{T-3}$.

(4) Finally, consider the case with three-sided private information (see Proposition 4.4). Since the first proposing stage is separating, if the proposer is rational, then the game proceeds as described by Proposition 4.3.1; otherwise, the game enters a continuation described by Proposition 4.3.2. Therefore, the probability that bargaining ends by period T is at least $(1 - \prod_{i \in N} \mu_i) [1 - (\frac{1}{3})]^{T-4}$.

4.5.2 Proof of Theorem 2

The result follows immediately from examining each ω -equivalence class of priors. Here I omit less interesting cases where exactly two agents have equal and positive prior probabilities of semi-rationality. WLOG consider priors such that $\mu_1 \leq \mu_2 \leq \mu_3$. (1) If $\mu^n \rightarrow \mathbf{0}$ is such that for all n , $\mu_1^n = \mu_2^n = \mu_3^n = 0$, then obviously $\lim_{n \rightarrow \infty} v^*(c, \mu^n) = (\frac{1}{3}, \frac{1}{3}, \frac{1}{3})$. (2) If

$\mu^n \rightarrow 0$ is such that for all n , $0 = \mu_1^n = \mu_2^n < \mu_3^n$, then

$$v_1^*(c, \mu^n) = v_2^*(c, \mu^n) = (1 - \mu_3^n) \frac{17}{54} + O(\mu_3^n)$$

and

$$v_3^*(c, \mu^n) = \frac{20}{54}.$$

Then we have $\lim_{n \rightarrow \infty} v^*(c, \mu^n) = (\frac{17}{54}, \frac{17}{54}, \frac{20}{54})$. Therefore, for n large enough,

$$v_1^*(c, \mu^n) = v_2^*(c, \mu^n) < v_3^*(c, \mu^n).$$

(3) If $\mu^n \rightarrow 0$ is such that for all n , $0 = \mu_1^n < \mu_2^n < \mu_3^n$, then

$$\begin{aligned} v_1^*(c, \mu^n) &= (1 - \mu_2^n)(1 - \mu_3^n) \frac{20}{81} + O(\mu_2^n) + O(\mu_3^n) \\ v_2^*(c, \mu^n) &= (1 - \mu_3^n) \frac{34}{81} + O(\mu_3^n) \\ v_3^*(c, \mu^n) &= (1 - \mu_2^n) \frac{27}{81} + O(\mu_2^n). \end{aligned}$$

And we have $\lim_{n \rightarrow \infty} v^*(c, \mu^n) = (\frac{20}{81}, \frac{34}{81}, \frac{27}{81})$. Therefore, for n large enough,

$$v_1^*(c, \mu^n) < v_3^*(c, \mu^n) < v_2^*(c, \mu^n).$$

(4). If $\mu^n \rightarrow 0$ is such that for all n , $0 < \mu_1^n < \mu_2^n < \mu_3^n$, then

$$\begin{aligned} v_1^*(c, \mu^n) &= (1 - \mu_2^n)(1 - \mu_3^n) \frac{34}{81} + O(\mu_2^n) + O(\mu_3^n) \\ v_2^*(c, \mu^n) &= (1 - \mu_3^n)(1 - \mu_1^n) \frac{27}{81} + O(\mu_3^n) + O(\mu_1^n) \\ v_3^*(c, \mu^n) &= (1 - \mu_1^n)(1 - \mu_2^n) \frac{20}{81} + O(\mu_1^n) + O(\mu_2^n). \end{aligned}$$

And we have $\lim_{n \rightarrow \infty} v^*(c, \mu^n) = (\frac{34}{81}, \frac{27}{81}, \frac{20}{81})$. Therefore, for n large enough,

$$v_3^*(c, \mu^n) < v_2^*(c, \mu^n) < v_1^*(c, \mu^n).$$

(5). If $\mu^n \rightarrow 0$ is such that for all n , $0 < \mu_1^n = \mu_2^n = \mu_3^n$, then for all $i \in N$,

$$v_i^*(c, \mu^n) = \prod_{j \neq i} (1 - \mu_j^n) \cdot \frac{1}{3} + \sum_{j \neq i} O(\mu_j^n).$$

And we have

$$\lim_{n \rightarrow \infty} v^*(c, \mu^n) = \left(\frac{1}{3}, \frac{1}{3}, \frac{1}{3} \right).$$

5 Conclusion and discussion

In this paper, I study the influence of reputation on majority bargaining in the simple three-agent divide-the-dollar environment with common conflicting claims on the share of surplus, a common discount factor and uniform recognition probabilities. When the rational types

are sufficiently patient or when bargaining interaction is sufficiently frequent, we construct a symmetric stationary perfect Bayesian equilibrium in which the bargaining process would reach a potential agreement in essentially finitely many periods. As time friction vanishes, efficiency loss vanishes. This Coasian feature is due to the fact that building up reputation incurs denouncement of semi-rationality and exclusion from winning coalitions. Although reputation does not have significant effect on efficiency of bargaining outcome, it has an interesting distributional implication. The rational type of the agent with the weakest reputation of commitment obtains the largest share of the surplus. Although strong reputation could be a curse, no reputation is not beneficial either. If the rational types could manipulate the strength of their reputations prior to the game, the resulting uncertainty would be infinitesimal.

There are many more interesting questions we can ask about the interaction between non-unanimity and reputation. The very first of them are on robustness of the features discovered in this specific equilibrium in this simple environment. Maintaining assumptions of common claim, common discount factor and uniform recognition probabilities, I anticipate that other symmetric stationary perfect Bayesian equilibria will exhibit the same feature. Relaxing the assumption of common claim, the strong efficiency result would fail. Consider $\frac{1}{2} < c_1 < c_2 < c_3$ and $0 < \mu_1 \leq \mu_2 < \mu_3 = 1$. It is easy to see that agent 1 and agent 2 would be involved in a war of attrition. However, the concealed rational type of agent 3 would actively reveal rationality, as otherwise he is excluded. It implies that in the case with three-sided private information, if agent 3 is rational, the game does not enter a war of attrition. Thus, when the prior probability of semi-rationality vanishes (after taking $\delta \rightarrow 1$ or $\Delta \rightarrow 0$), efficiency is restored. Introducing heterogeneous discount rates and departing from uniform recognition probabilities are also important, and it is not clear how these parameters would affect bargaining outcome. To fully answer the interesting and challenging questions above would call for more sophisticated approach than the constructive one adopted here.

6 Appendix

6.1 Proofs for one-sided private information

6.1.1 Proof of Proposition 4.2.1

Recall that

$$v^\delta(c, \mathbf{0}) = \left(\frac{1}{3}, \frac{1}{3}, \frac{1}{3}\right)$$

$$v^\delta(c, \mu^{\{3\}}) = \left(\frac{1}{3-\delta}, \frac{1}{3-\delta}, \frac{1}{3} - \frac{1}{9}\delta\right),$$

and thus

$$v^*(c, \mathbf{0}) = \left(\frac{1}{3}, \frac{1}{3}, \frac{1}{3}\right)$$

$$v^*(c, \mu^{\{3\}}) = \left(\frac{1}{2}, \frac{1}{2}, \frac{2}{9}\right).$$

Since $v_3^\delta(c, \mu^{\{3\}}) < c$, agent 3 can be screened.

Consider $\delta = 1$ and postulate that the proposal stage of agent 3 is separating. If the rational type of agent 3 reveals rationality, then he obtains $w_3^*(c, \mathbf{0}) = 1 - v^*(c, \mathbf{0})$; if he mimicks the semi-rational type, the game proceeds with posterior $\mu^{\{3\}}$ and he obtains $v_3^*(c, \mu^{\{3\}})$. Obviously $w_3^*(c, \mathbf{0}) > v_3^*(c, \mu^{\{3\}})$, and a separating proposal stage of agent 3 is indeed plausible. Now suppose both agent 1 and agent 2 screen agent 3 with the threat to enter a continuation with posterior $\mu^{\{3\}}$. Then the payoff from screening is $(1 - \mu_3) \left(1 - \frac{2}{9}\right) + \mu_3 \frac{1}{2}$. Since $c > \frac{1}{2}$ and the payoff from active concession to agent 3 is $1 - c < \frac{1}{2} < \frac{7}{9}$, active concession is dominated. The pre-recognition payoff for the rational type of agent $i \in \{1, 2\}$ at μ is

$$\begin{aligned} v_i^*(c, \mu) &= \frac{1}{3} \left[(1 - \mu_3) \left(1 - \frac{2}{9}\right) + \mu_3 \frac{1}{2} \right] \\ &\quad + \frac{1}{3} \left[(1 - \mu_3) 0 + \mu_3 \frac{1}{2} \right] \\ &\quad + \frac{1}{3} \left[(1 - \mu_3) \frac{1}{6} + \mu_3 \frac{1}{2} \right] \\ &= (1 - \mu_3) \frac{17}{54} + \mu_3 \frac{1}{2}. \end{aligned}$$

By symmetry, seeking agreement with agent $-i$ gives at most

$$\begin{aligned} 1 - v_i^*(c, \mu) &= (1 - \mu_3) \frac{37}{54} + \mu_3 \frac{1}{2} \\ &< (1 - \mu_3) \frac{7}{9} + \mu_3 \frac{1}{2}, \end{aligned}$$

and thus it is dominated. Since

$$\begin{aligned} v_3^*(c, \mu) &= \frac{1}{3} \cdot \frac{2}{9} + \frac{1}{3} \cdot \frac{2}{9} + \frac{1}{3} \cdot \frac{2}{3} \\ &> \frac{2}{9} \\ &= v_3^*(c, \mu^{\{3\}}), \end{aligned}$$

off-path response and belief updating can be specified as follows: if agent i proposes x with $x_3 < v_3^*(c, \mu^{\{3\}})$, then both types of agent 3 reject, and there is no belief updating; if agent i proposes x with $x_3 \in [v_3^*(c, \mu^{\{3\}}), c)$, the rational type of agent 3 accepts with probability one, while the semi-rational type rejects as prescribed; if agent i proposes x with $x_3 \geq c$, both types of agent 3 accept with probability one. As the relevant inequalities above are strict at $\delta = 1$, they are also satisfied when δ is close to 1. Therefore, the strategies described indeed consists an equilibrium.

6.1.2 Proof of Proposition 4.2.2

Recall that

$$v^\delta(c, \mu^{\{2,3\}}) = \left(1 - c, \frac{1}{2}c, \frac{1}{2}c\right)$$

and thus

$$v^*(c, \mu^{\{2,3\}}) = \left(1 - c, \frac{1}{2}c, \frac{1}{2}c\right).$$

Since $v_2^\delta(c, \mu^{\{2,3\}}) < c$, agent 2 can be screened with the threat to enter the continuation with posterior $\mu^{\{2,3\}}$. The payoff for agent 1 from screening agent 2 is then $(1 - \mu_2)(1 - \frac{1}{2}\delta c) + \mu_2\delta(1 - c)$, while by active concession, he obtains $1 - c$. Let $\bar{m}^\delta$ be uniquely defined by

$$(1 - \bar{m}^\delta) \left(1 - \frac{1}{2}\delta c\right) + \bar{m}^\delta \delta(1 - c) = 1 - c,$$

i.e., $\bar{m}^\delta$ is the cutoff belief above which screening is dominated by active concession. Obviously,

$$v_1^\delta(c, (0, \bar{m}^\delta, 1)) = 1 - c.$$

Let $\hat{m}^\delta$ be defined by

$$\begin{aligned} \frac{1}{\delta}(1 - c) &= \frac{1}{3} \left[(1 - \hat{m}^\delta) \left(1 - \frac{1}{2}\delta c\right) + \hat{m}^\delta \delta(1 - c) \right] \\ &\quad + \frac{2}{3}(1 - c) \end{aligned}$$

i.e., when the probability of agent 2 being semi-rational is $\hat{m}^\delta$, agent 1 is indifferent between passive concession and enter a continuation in which he would screen agent 2 if recognized but concede immediately otherwise. It can be shown that this strategy consists a perfect Bayesian equilibrium at $(0, \hat{m}^\delta, 1)$.

Note that $\lim_{\delta \rightarrow 1} \bar{m}^\delta = \lim_{\delta \rightarrow 1} \hat{m}^\delta = 1$. I claim that when δ is close to 1, $\hat{m}^\delta < \bar{m}^\delta$ and thus screening dominates active concession. Here we develop a trick which will be repeatedly used in proofs of later propositions. For $m \in (0, 1)$, let $l^\delta(m) = \frac{1}{(1-\delta)(1-c)} \frac{1-m}{m}$. Then for all $\delta \in (0, 1)$, $m > m'$ if and only if $l^\delta(m) < l^\delta(m')$. From the definition of $\bar{m}^\delta$, we have

$$l^\delta(\bar{m}^\delta) \left(1 - \frac{1}{2}\delta c\right) + \frac{\delta(1-c)}{(1-\delta)(1-c)} = l^\delta(\bar{m}^\delta)(1-c) + \frac{1-c}{(1-\delta)(1-c)},$$

and thus

$$l^\delta(\bar{m}^\delta) = \frac{1}{c - \frac{1}{2}\delta c}.$$

Let

$$\bar{l} = \lim_{\delta \rightarrow 1} l^\delta(\bar{m}^\delta) = \frac{2}{c}.$$

By a similar operation on $\hat{m}^\delta$, we have

$$\hat{l} = \lim_{\delta \rightarrow 1} l^\delta(\hat{m}^\delta) = \frac{8}{c}.$$

Since $\hat{l} > \bar{l}$, for δ close to 1, we have $l^\delta(\hat{m}^\delta) > l^\delta(\bar{m}^\delta)$ and $\hat{m}^\delta < \bar{m}^\delta$.

To see that the rational types agent 2 and agent 3 would indeed force concession by agent 1, by Proposition 4.1.2 and Proposition 4.2.1, they obtain approximately $\frac{1}{2}$ by revealing rationality when δ is close to 1, but they can obtain c by forcing concession. Given the strategy of agent 1 at $(0, \hat{m}^\delta, 1)$, the pre-recognition payoff for the rational type of agent 3

at $(0, \hat{m}^\delta, 1)$ is

$$\begin{aligned} v_3^\delta(c, (0, \hat{m}^\delta, 1)) &= \frac{1}{3} \left[(1 - \hat{m}^\delta) \cdot 0 + \hat{m}^\delta \delta v_3^\delta(c, \mu^{\{2,3\}}) \right] + \\ &\quad \frac{1}{3} \cdot 0 + \frac{1}{3} c \\ &= \left(\frac{1}{3} + \frac{1}{6} \hat{m}^\delta \delta \right) c. \end{aligned}$$

The pre-recognition payoff for the rational type of agent 2 at $(0, \hat{m}^\delta, 1)$ is

$$\begin{aligned} v_2^\delta(c, (0, \hat{m}^\delta, 1)) &= \frac{1}{3} \cdot \frac{1}{2} \delta c + \frac{1}{3} c + \frac{1}{3} \cdot 0 \\ &= \left(\frac{1}{3} + \frac{1}{6} \delta \right) c \\ &< v_2^\delta(c, \mu^{\{2,3\}}), \end{aligned}$$

which implies that if agent 1 proposes x with $x_2 \in (\delta v_2^\delta(c, (0, \hat{m}^\delta, 1)), \delta v_2^\delta(c, \mu^{\{2,3\}}))$, the voting decision of the rational type of agent 2 cannot be pure. Postulate that for any x with $x_2 \in [\delta v_2^\delta(c, (0, \hat{m}^\delta, 1)), \delta v_2^\delta(c, \mu^{\{2,3\}})]$, the rational type of agent 2 mixes between acceptance and rejection. By rejection of x_2 , the game proceed with posterior $(0, \bar{m}^\delta, 1)$. If agent 1 is recognized at $(0, \bar{m}^\delta, 1)$, he mixes between screening and concession such that the pre-recognition payoff for the rational type of agent 2 at $(0, \bar{m}^\delta, 1)$ is $\frac{1}{6} x_2$. Then the payoff for agent 1 from proposing x with $x_2 \in [\delta v_2^\delta(c, (0, \hat{m}^\delta, 1)), \delta v_2^\delta(c, \mu^{\{2,3\}})]$ and $x_1 = 1 - x_2$ is

$$\begin{aligned} &\left(1 - \frac{\hat{m}^\delta}{\bar{m}^\delta} \right) (1 - x_2) + \frac{\hat{m}^\delta}{\bar{m}^\delta} \delta (1 - c) \\ &\leq \left(1 - \frac{\hat{m}^\delta}{\bar{m}^\delta} \right) (1 - \delta v_2^\delta(c, (0, \hat{m}^\delta, 1))) + \frac{\hat{m}^\delta}{\bar{m}^\delta} \delta (1 - c). \end{aligned}$$

Let $\tilde{m}^\delta$ be defined by

$$\left(1 - \frac{\tilde{m}^\delta}{\bar{m}^\delta} \right) \left[1 - \delta \left(\frac{1}{3} + \frac{1}{6} \delta \right) c \right] + \frac{\tilde{m}^\delta}{\bar{m}^\delta} \delta (1 - c) = (1 - \tilde{m}^\delta) \left(1 - \frac{1}{2} \delta c \right) + \tilde{m}^\delta \delta (1 - c)$$

i.e., $\tilde{m}^\delta$ is the cutoff above which agent 1 prefers offering $\delta v_2^\delta(c, \mu^{\{2,3\}})$ to agent 2 to offering $\delta v_2^\delta(c, (0, \hat{m}^\delta, 1))$. To justify the strategy of agent 1 at $(0, \hat{m}^\delta, 1)$, it remains to show that $\hat{m}^\delta > \tilde{m}^\delta$. Since

$$l^\delta(\hat{m}^\delta) \cdot \frac{1}{6} (1 - \delta) \delta c = l^\delta(\bar{m}^\delta) \cdot \left[\delta (1 - c) - \left(1 - \frac{1}{2} \delta c \right) \right],$$

we have

$$\lim_{\delta \rightarrow 1} l^\delta(\tilde{m}^\delta) = \infty.$$

And thus for δ close to 1, we have $l^\delta(\hat{m}^\delta) < l^\delta(\tilde{m}^\delta)$ and $\hat{m}^\delta > \tilde{m}^\delta$. To support a semi-pooling proposal stage of agent 2 at $(0, \mu_2, 1)$, let the posterior updated from observing $(1 - c, c, 0)$ be $\hat{m}^\delta$ and let $r^\delta(c, \mu)$ be defined by

$$1 - \delta v_1^\delta(c, \mu^{\{3\}}) = (1 - r^\delta) c + r^\delta \delta v_2^\delta(c, (0, \hat{m}^\delta, 1)).$$

Let $\delta \rightarrow 1$, we also have

$$r^* = 2 - \frac{1}{c}.$$

Consider δ close to 1 such that $\mu_2 < \hat{m}^\delta$. We next construct the equilibrium play at μ . If $c > \frac{2}{3}$, then for δ close to $1 - c \leq \frac{1}{2}\delta c$ and $(0, 1 - c, c)$ is not a screening proposal. Postulate that neither of agent 1 and agent 2 passively concede to agent 3. To justify this, let

$$\begin{aligned} v_1^\delta(c, \mu) &= \frac{1}{3} \left[(1 - \mu_2) \left(1 - \delta v_2^\delta(c, \mu^{\{2,3\}}) \right) + \mu_2 \delta v_1^\delta(c, \mu^{\{2,3\}}) \right] \\ &\quad + \frac{1}{3} \left[\left(1 - \frac{\mu_2}{\hat{m}^\delta} \right) \delta v_1^\delta(c, \mu^{\{3\}}) + \frac{\mu_2}{\hat{m}^\delta} (1 - c) \right] \\ &\quad + \frac{1}{3} \delta v_1^\delta(c, \mu) \end{aligned}$$

and

$$v_2^\delta(c, \mu) = \frac{1}{3} \delta v_2^\delta(c, \mu^{\{2,3\}}) + \frac{1}{3} \left(1 - \delta v_1^\delta(c, \mu^{\{3\}}) \right) + \frac{1}{3} \delta v_2^\delta(c, \mu).$$

Let $\delta \rightarrow 1$, we have

$$\begin{aligned} v_1^*(c, \mu) &= (1 - \mu_2) \left(\frac{3}{4} - \frac{1}{4}c \right) + \mu_2 (1 - c) \\ v_2^*(c, \mu) &= \frac{1}{4}c + \frac{1}{4}. \end{aligned}$$

Obviously, when $c > \frac{2}{3}$ and $\mu_2 < 1$, $v_1^*(c, \mu) > 1 - c$ and $v_2^*(c, \mu) > 1 - c$. Moreover, $v_2^*(c, \mu^{\{2,3\}}) < v_2^*(c, \mu)$, implying that we can set any proposal x by agent 1 with $x_2 < v_2^*(c, \mu^{\{2,3\}})$ to be rejected by both types of agent 2. Since the payoff from screening agent 2 at $\delta = 1$ is

$$(1 - \mu_2) \left(1 - \frac{1}{2}c \right) + \mu_2 (1 - c) > v_1^*(c, \mu),$$

agent 1 does not make non-serious proposals when δ is close to 1. Then we show that the concealed rational type of agent 3 would actively reveal rationality at μ . Suppose it is optimal for him to maintain concealed at μ , we have

$$\begin{aligned} v_3^\delta(c, \mu) &= \frac{1}{3} \mu_2 \delta v_3^\delta(c, \mu^{\{2,3\}}) \\ &\quad + \frac{1}{3} \frac{\mu_2}{\hat{m}^\delta} r^\delta \delta v_3^\delta(c, (0, \hat{m}^\delta, 1)) \\ &\quad + \frac{1}{3} \delta v_3^\delta(c, \mu). \end{aligned}$$

Let $\delta \rightarrow 1$, we have

$$v_3^*(c, \mu) = \mu_2 \left(\frac{3}{4}c - \frac{1}{4} \right)$$

However, if he actively reveals rationality and plays according to posterior $proj_{1,3;\emptyset}\mu$, he obtains

$$(1 - \mu_2) \frac{7}{9} + \mu_2 \frac{1}{2} > \mu_2 \left(\frac{3}{4}c - \frac{1}{4} \right),$$

which is a contradiction. Therefore,

$$\begin{aligned} v_3^\delta(c, \mu) &= \frac{1}{3} \mu_2 \delta v_3^\delta(c, \mu^{\{2,3\}}) \\ &\quad + \frac{1}{3} \frac{\mu_2}{\hat{m}^\delta} r^\delta \delta v_3^\delta(c, (0, \hat{m}^\delta, 1)) \\ &\quad + \frac{1}{3} w_3^\delta(3, c, \text{proj}_{1,3;\emptyset} \mu) \end{aligned}$$

and

$$v_3^*(c, \mu) = (1 - \mu_2) \frac{7}{27} + \mu_2 \frac{1}{2} c$$

If $c \leq \frac{2}{3}$ then $1 - c > \frac{1}{2} \delta c$ and $(0, 1 - c, c)$ is a screening proposal. It is easy to see that if agent 3 only seeks agreement with agent 1, agent 1 does not concede to agent 3. Suppose agent 3 only forces concession by agent 2. Then

$$\begin{aligned} v_1^*(c, \mu) &= (1 - \mu_2) \left(\frac{1}{2} - \frac{1}{6} c \right) + \mu_2 (1 - c) \\ v_2^*(c, \mu) &= \frac{1}{2} - \frac{1}{6} c. \end{aligned}$$

As semi-rational types maximize probability of acceptance of their claims, it is necessary that $v_1^*(c, \mu) \geq 1 - c$, or $c \geq \frac{3}{5}$. Since $v_2^*(c, \mu) > v_2^*(c, \mu^{\{2,3\}})$, we can set any proposal x by agent 1 with $x_2 < v_2^*(c, \mu^{\{2,3\}})$ to be rejected by both types of agent 2. If $c < \frac{3}{5}$, agent 3 needs to mix between forcing concession by agent 1 and agent 2. To balance the incentive of agent 3, agent 1 passively concedes with probability μ_2 . Let $p^\delta(c, \mu)$ be defined by

$$\begin{aligned} v_1^\delta(c, \mu) &= \frac{1 - c}{\delta} \\ &= \frac{1}{3} \left[(1 - \mu_2) \left(1 - \delta v_2^\delta(c, \mu^{\{2,3\}}) \right) + \mu_2 \delta v_1^\delta(c, \mu^{\{2,3\}}) \right] \\ &\quad + \frac{1}{3} \left[\left(1 - \frac{\mu_2}{\hat{m}^\delta} \right) \delta v_1^\delta(c, \mu^{\{3\}}) + \frac{\mu_2}{\hat{m}^\delta} (1 - c) \right] \\ &\quad + \frac{1}{3} \left[p^\delta (1 - c) + (1 - p^\delta) \mu_2 \delta v_1^\delta(c, \mu^{\{2,3\}}) \right]. \end{aligned}$$

Let $\delta \rightarrow 1$, we have

$$p^* = \frac{3 - 5c}{2(1 - c)}$$

and

$$\begin{aligned} v_1^*(c, \mu) &= 1 - c \\ v_2^*(c, \mu) &= \frac{2c}{3 - p\mu_2} \\ &> v_2^*(c, \mu^{\{2,3\}}). \end{aligned}$$

Then we can set any proposal x by agent 1 with $x_2 < v_2^*(c, \mu^{\{2,3\}})$ to be rejected by both types of agent 2. It is easy to see that when $c \leq \frac{2}{3}$, $v_1^*(c, \mu) < w_1^*(1, c, \mu)$ and thus agent 1 does not make non-serious offers. Finally, it can be shown that the concealed rational

type of agent 3 would actively reveal rationality. Suppose it is optimal for him to maintain concealed at μ , then

$$w_3^*(3, c, \mu) = \max \left\{ (1 - \mu_2) c + \mu_2 v_3^*(c, \mu), (1 - \mu_2) c + \mu_2 v_3^*(c, \mu^{\{2,3\}}) \right\}$$

$$v_3^*(c, \mu) = \mu_2 \left(\frac{1}{2} c - \frac{1}{6} \right) + \frac{1}{3} w_3^*(3, c, \mu),$$

which yields

$$w_3^*(3, c, \mu) = (1 - \mu_2) c + \mu_2 \frac{1}{2} c.$$

However, by actively revealing rationality, he obtains

$$(1 - \mu_2) \frac{7}{9} + \mu_2 \frac{1}{2} > w_3^*(3, c, \mu),$$

which is a contradiction. Therefore,

$$\begin{aligned} v_3^\delta(c, \mu) &= \frac{1}{3} \mu_2 \delta v_3^\delta(c, \mu^{\{2,3\}}) \\ &\quad + \frac{1}{3} \frac{\mu_2}{\hat{m}^\delta} r^\delta \delta v_3^\delta(c, (0, \hat{m}^\delta, 1)) \\ &\quad + \frac{1}{3} w_3^\delta(3, c, proj_{1,3;\emptyset} \mu) \end{aligned}$$

and

$$v_3^*(c, \mu) = (1 - \mu_2) \frac{7}{27} + \mu_2 \frac{1}{2} c.$$

6.2 Proofs for Section 4.3

6.2.1 Proof of Proposition 4.3.1

WLOG, let $\mu_2 \leq \mu_3$. I first show that for agent 1, when δ is close to 1, the payoff from screening agent 2 by the threat to enter the continuation with posterior $proj_{1;2} \mu$ is weakly higher than the payoff from screening agent 3 by the threat of posterior $proj_{1;3} \mu$. From Proposition 4.2.2, we know that

$$\begin{aligned} v_2^*(c, proj_{1;2} \mu) &= (1 - \mu_3) \frac{7}{27} + \mu_3 \frac{1}{2} c \\ v_3^*(c, proj_{1;3} \mu) &= (1 - \mu_2) \frac{7}{27} + \mu_2 \frac{1}{2} c, \end{aligned}$$

And for each $i \in \{2, 3\}$,

$$v_1^*(c, proj_{1;i} \mu) = (1 - \mu_{-i}) [(1 - c) + f(c)] + \mu_{-i} (1 - c)$$

in which

$$\begin{aligned} f(c) &= 0, & \text{if } c < \frac{3}{5} \\ &= \frac{5}{6} c - \frac{1}{2}, & \text{if } c \in \left[\frac{3}{5}, \frac{2}{3} \right] \\ &= \frac{3}{4} c - \frac{1}{4}, & \text{if } c > \frac{2}{3}. \end{aligned}$$

At $\delta = 1$, let the payoff for agent 1 from screening agent i be denoted by

$$s^*(i, c, \mu) = (1 - \mu_i) (1 - v_i^*(c, proj_{1;i} \mu)) + \mu_i v_1^*(c, proj_{1;i} \mu)$$

Then

$$\begin{aligned}
s^*(i, c, \mu) - s^*(-i, c, \mu) &= \left[\mu_i \frac{1}{2} c + \mu_{-i} (c - f(c)) \right] \\
&\quad - \left[\mu_{-i} \frac{1}{2} c + \mu_i (c - f(c)) \right] \\
&= (\mu_{-i} - \mu_i) \left(\frac{1}{2} c - f(c) \right).
\end{aligned}$$

It is easy to see that $\frac{1}{2}c - f(c) > 0$, and thus $s^*(i, c, \mu) > s^*(-i, c, \mu)$ if and only if $\mu_{-i} > \mu_i$. Therefore, when δ is close to 1, agent 1 prefers to screen agent 2. We also know from Proposition 4.2.2 that for all $i \in \{2, 3\}$, if agent i is recognized first, it is plausible that his proposal stage is separating: by mimicking the semi-rational type, the game proceeds with posterior $proj_{1;i}\mu$, which gives lower payoff than revealing rationality and screening agent $-i$. Given the strategy described, denote the corresponding pre-recognition payoffs for rational types by $v^*(c, \mu)$. It remains to show that (1) $v_i^*(c, proj_{1;i}\mu) < v_i^*(c, \mu)$ and thus we can set any proposal x by agent 1 with $x_2 < v_i^*(c, proj_{1;i}\mu)$ to be rejected by both types of agent i , and (2) $s^*(2, c, \mu) > v_1^*(c, \mu)$ and thus agent 1 does not make non-serious offers. I omit the detailed calculation here.

6.2.2 Proof of Proposition 4.3.2

First consider $c \leq \frac{2}{3}$. For all $i \in \{2, 3\}$, since

$$\begin{aligned}
w_i^\delta(i, c, proj_{1;i}\mu) &= (1 - \mu_{-i}) \left(1 - \frac{1}{2}\delta c \right) + \mu_{-i}\delta(1 - c) \\
&> (1 - \mu_{-i})c + \mu_{-i}\delta(1 - c),
\end{aligned}$$

a separating proposal stage of agent i at μ is plausible. Postulating separation, I claim that when δ is close to 1, there is an equilibrium in which neither of agent 2 and agent 3 passively concedes to agent 1 at μ . To see this, given no passive concession, we have that for all $i \in \{2, 3\}$,

$$\begin{aligned}
v_i^*(c, \mu) &= \frac{1}{3} \left[(1 - \mu_{-i}) \left(1 - \frac{1}{2}c \right) + \mu_{-i}(1 - c) \right] \\
&\quad + \frac{1}{3} \left[(1 - \mu_{-i}) \frac{1}{2}c + \mu_{-i}(1 - c) \right] \\
&\quad + \frac{1}{3} v_i^*(c, \mu) \\
&= (1 - \mu_{-i}) \frac{1}{2} + \mu_{-i}(1 - c) \\
&> 1 - c.
\end{aligned}$$

Next it can be shown that the concealed rational type of agent 1 would actively revealed rationality. Suppose not, then by no passive concession,

$$\begin{aligned}
w_1^*(1, c, \mu) &= v_1^*(c, \mu) \\
&= \sum_{i \in \{2, 3\}} \frac{1}{3} \left[(1 - \mu_i) \mu_{-i} \frac{1}{2} c + \mu_i \mu_{-i} (1 - c) \right] \\
&\quad + \frac{1}{3} v_1^*(c, \mu) \\
&= (1 - \mu_2) \mu_3 \frac{1}{2} c + \mu_2 (1 - \mu_3) \frac{1}{2} c + \mu_2 \mu_3 (1 - c)
\end{aligned}$$

However, if he reveals rationality, by Proposition 4.3.1, he obtains

$$\begin{aligned}
w_1^*(1, c, proj_{1; \emptyset} \mu) &\geq (1 - \mu_2) (1 - \mu_3) \frac{20}{27} \\
&\quad + (1 - \mu_2) \mu_3 \left(1 - \frac{1}{2} c \right) \\
&\quad + \mu_2 (1 - \mu_3) [(1 - c) + f(c)] \\
&\quad + \mu_2 \mu_3 (1 - c)
\end{aligned}$$

in which the RHS of the inequality is the payoff from screening agent 2. Obviously, $w_1^*(1, c, proj_{1; \emptyset} \mu) > w_1^*(1, c, \mu)$ which is a contradiction.

Second consider $c > \frac{2}{3}$. For all $i \in \{2, 3\}$, since when δ is close to 1,

$$w_i^\delta(i, c, proj_{i; 1} \mu) < (1 - \mu_{-i}) c + \mu_{-i} \delta (1 - c)$$

a separating proposal stage of agent i at μ is implausible. Suppose agent 2 is recognized in period 0. The semi-rational type of agent 2 proposes $(0, c, 1 - c)$. Observation of this proposal induces posterior $(1, \hat{m}^{\delta, (1)}, \mu_3)$, with $\hat{m}^{\delta, (1)}$ defined by

$$\frac{1}{\delta} (1 - c) = \frac{1}{3} w_3^\delta \left(3, c, \left(1, \hat{m}^{\delta, (1)}, 0 \right) \right) + \frac{2}{3} (1 - c)$$

i.e., the rational type of agent 3 is indifferent between acceptance and rejection if he plans to actively reveal rationality and passively concede in period 1. Then $\hat{m}^{\delta, (1)}$ is equal to $\hat{m}^\delta$ defined in the proof of Proposition 4.2.2, and thus

$$\tilde{l}^{(1)} = \lim_{\delta \rightarrow 1} l^\delta \left(\hat{m}^{\delta, (1)} \right) = \frac{8}{c}.$$

By rejection, let posterior be updated to $(1, \hat{m}^{\delta, (1)}, \hat{m}^{\delta, (2)})$. Suppose agent 1 or agent 2 is recognized in period 1. Postulating that the rational type of agent 3 passively concede, the voting stage of agent 3 is separating. Based on Proposition 4.2.3, let $\tilde{m}^{\delta, (1)}$ be defined by

$$\frac{1}{\delta} (1 - c) = v_3^\delta \left(c, \left(1, \tilde{m}^{\delta, (1)}, 1 \right) \right),$$

i.e., $\tilde{m}^{\delta, (1)}$ is the cutoff above which a separating voting stage of agent 3 is plausible. Then we have

$$\tilde{l}^{(1)} = \lim_{\delta \rightarrow 1} l^\delta \left(\tilde{m}^{\delta, (1)} \right) = \frac{8}{3c - 2}.$$

Since $\tilde{l}^{(1)} < \tilde{l}^{(1)}$, when δ is close to 1, we have $l^\delta(\hat{m}^{\delta,(1)}) < l^\delta(\tilde{m}^{\delta,(1)})$, and thus $\hat{m}^{\delta,(1)} > \tilde{m}^{\delta,(1)}$ as desired. Suppose agent 3 is recognized in period 2. The semi-rational type of agent 3 proposes $(0, 1 - c, c)$. Observation of this proposal induces posterior $(1, \hat{m}^{\delta,(1)}, \bar{m}^{\delta,(2)})$, with $\bar{m}^{\delta,(2)}$ be defined by

$$\begin{aligned} \frac{1}{\delta}(1 - c) &= \frac{1}{3} \left[\left(1 - \bar{m}^{\delta,(2)}\right) c + \bar{m}^{\delta,(2)} \delta (1 - c) \right] \\ &+ \frac{1}{3} (1 - c) + \frac{1}{3} \bar{m}^{\delta,(2)} \delta (1 - c), \end{aligned}$$

i.e., the rational type of agent 2 is indifferent between acceptance and rejection of $1 - c$, if in period 3, agent 2 and agent 3 force concession by one another, and agent 1 force concession by agent 3. Then we have

$$l^{(2)} = \lim_{\delta \rightarrow 1} l^\delta(\bar{m}^{\delta,(2)}) = \frac{5}{3c - 2}.$$

By rejection, let posterior be updated to $(1, \bar{m}^{\delta,(1)}, \bar{m}^{\delta,(2)})$. Let

$$v_3^\delta \left(c, \left(1, \bar{m}^{\delta,(1)}, \bar{m}^{\delta,(2)}\right) \right) = \frac{1}{3} \left[\left(1 - \bar{m}^{\delta,(1)}\right) c + \bar{m}^{\delta,(1)} \delta (1 - c) \right] + \frac{2}{3} (1 - c)$$

be the pre-recognition payoff for the rational type of agent 3 at $(1, \bar{m}^{\delta,(1)}, \bar{m}^{\delta,(2)})$, if he actively forces concession by agent 2 and passively concedes. To support a semi-pooling proposal stage of agent 3 at $(1, \hat{m}^{\delta,(1)}, \hat{m}^{\delta,(2)})$, let $\bar{m}^{\delta,(1)}$ be defined by

$$w_3^\delta \left(3, c, \left(1, \hat{m}^{\delta,(1)}, 0\right) \right) = \left(1 - \frac{\hat{m}^{\delta,(1)}}{\bar{m}^{\delta,(1)}}\right) c + \frac{\hat{m}^{\delta,(1)}}{\bar{m}^{\delta,(1)}} \delta v_3^\delta \left(c, \left(1, \bar{m}^{\delta,(1)}, \bar{m}^{\delta,(2)}\right) \right).$$

It can be shown that

$$\bar{l}^{(1)} = \lim_{\delta \rightarrow 1} l^\delta(\bar{m}^{\delta,(1)}) = \frac{19 - \frac{12}{c}}{2c - 1}.$$

Postulating that the rational types of agent 2 and agent 3 passively concedes, the corresponding voting stages are separating. Since $\max \left\{ \bar{l}^{(1)}, \bar{l}^{(2)} \right\} < \tilde{l}^{(1)}$, we have when δ is close to 1, $\min \left\{ \bar{m}^{\delta,(1)}, \bar{m}^{\delta,(2)} \right\} > \tilde{m}^{\delta,(1)}$ and thus separation is indeed plausible. Note that existence of the cutoffs defined above relies on δ being close to 1, and for our construction to be meaningful, μ_2 and μ_3 need to be lower than these cutoffs. Then the payoff for the rational type of agent 2 at $(1, \hat{m}^{\delta,(1)}, \hat{m}^{\delta,(2)})$ with strategies defined above can be expressed as

$$\begin{aligned} v_2^\delta \left(c, \left(1, \hat{m}^{\delta,(1)}, \hat{m}^{\delta,(2)}\right) \right) &= \frac{1}{3} \left[\left(1 - \hat{m}^{\delta,(2)}\right) c + \hat{m}^{\delta,(2)} \delta (1 - c) \right] \\ &+ \frac{1}{3} \left[\left(1 - \frac{\hat{m}^{\delta,(2)}}{\bar{m}^{\delta,(2)}}\right) \frac{1}{2} \delta c + \frac{\hat{m}^{\delta,(2)}}{\bar{m}^{\delta,(2)}} (1 - c) \right] \\ &+ \frac{1}{3} \hat{m}^{\delta,(2)} \delta (1 - c). \end{aligned}$$

To support a semi-pooling proposal stage of agent 2 at μ , let $\hat{m}^{\delta,(2)}$ be defined by

$$w_3^\delta \left(3, c, (1, 0, \mu_3) \right) = \left(1 - \frac{\mu_2}{\hat{m}^{\delta,(2)}}\right) c + \frac{\mu_2}{\hat{m}^{\delta,(2)}} \delta v_2^\delta \left(c, \left(1, \hat{m}^{\delta,(1)}, \hat{m}^{\delta,(2)}\right) \right).$$

It can be shown that

$$\lim_{\delta \rightarrow 1} \frac{1 - \hat{m}^{\delta, (2)}}{\hat{m}^{\delta, (2)}} = \frac{1 - \mu_2}{\mu_2} \left(3 - \frac{2}{c} \right).$$

Then for δ close to 1, there is an equilibrium in which neither agent 2 and agent 3 passively concede to agent 1 at μ : let $\delta \rightarrow 1$, for all $i \in \{2, 3\}$, we have

$$\begin{aligned} v_i^*(c, \mu) &= \frac{1}{3} v_i^*(c, \mu) \\ &\quad + \frac{1}{3} \left[(1 - \mu_{-i}) \left(1 - \frac{1}{2}c \right) + \mu_{-i} (1 - c) \right] \\ &\quad + \frac{1}{3} \left[(1 - \mu_{-i}) \frac{1}{2}c + \mu_{-i} (1 - c) \right] \\ &= (1 - \mu_{-i}) \frac{1}{2} + \mu_{-i} (1 - c) \\ &> 1 - c. \end{aligned}$$

If agent 3 recognized first at μ , let strategies be defined symmetrically. For the concealed rational type of agent 1, I claim that he would not actively reveal rationality at $(1, \hat{m}^{\delta, (1)}, \hat{m}^{\delta, (2)})$ when δ is close to 1. To see this, note that $\hat{m}^{*, (2)} = \lim_{\delta \rightarrow 1} \hat{m}^{\delta, (2)} < \lim_{\delta \rightarrow 1} \hat{m}^{\delta, (1)} = 1$, and thus the payoff from revealing rationality and screening agent 2 is approximately

$$\left(1 - \hat{m}^{*, (2)} \right) \left(1 - \frac{1}{2}c \right) + \hat{m}^{*, (2)} (1 - c).$$

Given $c > \frac{2}{3}$, it is lower than the payoff from forcing concession by agent 2

$$\left(1 - \hat{m}^{*, (2)} \right) c + \hat{m}^{*, (2)} (1 - c)$$

Since $\lim_{\delta \rightarrow 1} \bar{m}^{\delta, (1)} = \lim_{\delta \rightarrow 1} \bar{m}^{\delta, (2)} = 1$, whether the concealed rational type of agent 1 actively reveal rationality at $(1, \bar{m}^{\delta, (1)}, \bar{m}^{\delta, (2)})$ does not make a difference in terms of payoff when δ is close to 1. Finally, it can be shown that the concealed rational type of agent 1 actively reveals rationality at μ . Let

$$\begin{aligned} v_1^*(c, (1, 1, \hat{m}^{*, (2)})) &= \frac{1}{3} \left[\left(1 - \hat{m}^{*, (2)} \right) c + \hat{m}^{*, (2)} (1 - c) \right] \\ &\quad + \frac{1}{3} \hat{m}^{*, (2)} (1 - c) \\ &\quad + \frac{1}{3} \left[\left(1 - \hat{m}^{*, (2)} \right) \frac{1}{2}c + \hat{m}^{*, (2)} (1 - c) \right] \\ &= \left(1 - \hat{m}^{*, (2)} \right) \frac{1}{2}c + \hat{m}^{*, (2)} (1 - c). \end{aligned}$$

be the limit of his payoff at $(1, \hat{m}^{\delta, (1)}, \hat{m}^{\delta, (2)})$ as $\delta \rightarrow 1$. Suppose he maintains concealed at

μ , then

$$\begin{aligned} v_1^*(c, \mu) &= \frac{1}{3} v_1^*(c, \mu) \\ &+ \frac{1}{3} \left[(1 - \mu_2) \mu_3 \frac{1}{2} c + \frac{\mu_2 \mu_3}{\hat{m}^{*,(2)}} v_1^* \left(c, \left(1, 1, \hat{m}^{*,(2)} \right) \right) \right] \\ &+ \frac{1}{3} \left[(1 - \mu_3) \mu_2 \frac{1}{2} c + \frac{\mu_2 \mu_3}{\hat{m}^{*,(2')}} v_1^* \left(c, \left(1, \hat{m}^{*,(2')}, 1 \right) \right) \right] \end{aligned}$$

in which $\hat{m}^{*,(2')}$ is defined by

$$\frac{1 - \hat{m}^{*,(2')}}{\hat{m}^{*,(2')}} = \frac{1 - \mu_3}{\mu_3} \left(3 - \frac{2}{c} \right).$$

And we have

$$\begin{aligned} v_1^*(c, \mu) &= (1 - \mu_2) \mu_3 \cdot \left(c - \frac{1}{2} \right) \\ &+ (1 - \mu_3) \mu_2 \cdot \left(c - \frac{1}{2} \right) \\ &+ \mu_1 \mu_2 (1 - c) \\ &< w_1^*(1, c, proj_{1;\emptyset} \mu), \end{aligned}$$

which is a contradiction.

References

- [1] Abreu, D., & Gul, F. (2000). Bargaining and reputation. *Econometrica*, 68(1), 85-117.
- [2] Ali, S. N. M. (2006). Waiting to settle: Multilateral bargaining with subjective biases. *Journal of Economic Theory*, 130(1), 109-137.
- [3] Ausubel, L. M., Cramton, P., & Deneckere, R. J. (2002). Bargaining with incomplete information. *Handbook of game theory with economic applications*, 3, 1897-1945.
- [4] Banks, J. S., & Duggan, J. (2000). A bargaining model of collective choice. *The American Political Science Review*, 94(1), 73.
- [5] Banks, J. S., & Duggan, J. (2006). A general bargaining model of legislative policy-making. *Quarterly Journal of Political Science*, 1(1), 49-85.
- [6] Baron, D. P., & Ferejohn, J. A. (1989). Bargaining in legislatures. *American political science review*, 83(04), 1181-1206.
- [7] Chen, Jidong, Sequential Agenda Setting with Strategic and Informative Voting (April 21, 2017). Available at SSRN: <https://ssrn.com/abstract=2429028> or <http://dx.doi.org/10.2139/ssrn.2429028>

- [8] Chen, Y., & Eraslan, H. (2014). Rhetoric in legislative bargaining with asymmetric information. *Theoretical Economics*, 9(2), 483-513.
- [9] Eraslan, H. (2002). Uniqueness of stationary equilibrium payoffs in the Baron–Ferejohn model. *Journal of Economic Theory*, 103(1), 11-30.
- [10] Eraslan, H., & Merlo, A. (2002). Majority rule in a stochastic model of bargaining. *Journal of Economic Theory*, 103(1), 31-48.
- [11] Kambe, S. (1999). Bargaining with imperfect commitment. *Games and Economic Behavior*, 28(2), 217-237.
- [12] Tsai, T. S., & Yang, C. C. (2010). On majoritarian bargaining with incomplete information. *International Economic Review*, 51(4), 959-979.