

Communication in Bargaining Over Decision Rights

Wooyoung Lim *

Department of Economics

The Hong Kong University of Science and Technology

November 25, 2013

Abstract

This paper develops a model of bargaining over decision rights between an uninformed principal and an informed but self-interested agent. We introduce two different bargaining mechanisms: tacit and explicit bargaining. In tacit bargaining, an uninformed principal makes a take-it-or-leave-it price offer to the agent, who then decides whether to accept or reject the offer. In the equilibrium of the game, the principal inefficiently screens out some agent types so that the agent's private information cannot be fully utilized when the decision is made. In explicit bargaining in which parties can communicate explicitly via cheap talk before tacit bargaining, however, an equilibrium with no such inefficient screening exists even when the conflict of interest is arbitrarily large. We also follow a mechanism design approach, showing that under certain conditions, explicit bargaining is an optimal bargaining mechanism that maximizes the joint surplus of the parties.

Keywords: decision rights, delegation, monetary transfer, cheap talk.

JEL classification: D23, D83, L24.

*This paper is based on the second chapter of my Ph.D. dissertation, which has been submitted to the University of Pittsburgh. I thank two anonymous referees and an associate editor for their detailed comments and helpful suggestions. I am grateful to Andreas Blume and Oliver Board for their invaluable guidance, encouragement and support. I also would like to thank Sourav Bhattacharya, Yeon-Koo Che, Won-Jea Huh, Maxim Ivanov, Navin Kartik, Jonathan Lafky, Ernest K. Lai, Jinghong P. Liang, Yuanchuan Lien, Tymofiy Mylovanov, Jay Schwarz, Joel Sobel, Vasiliki Skreta, Yeolyong Sung, and Gábor Virág for their helpful discussions and remarks. Seminar participants at HKU, HKUST, PSU, SKKU, Sogang University, University of Pittsburgh, University of Bonn, and conference participants at the 2009 North American Summer Meeting of the Econometric Society, 2009 ECORE Summer School, 20th Stony Brook Game Theory Festival and 2009 Spring and Fall Midwest Theory Meeting, 2013 ASSA Meeting, and 2013 SAET Meeting contributed valuable comments and insights. In addition, I am indebted to Maria Goltsman for a series of discussions and comments that greatly improved this paper. The financial support of the Andrew Mellon Pre-Doctoral Fellowship at the University of Pittsburgh is also gratefully acknowledged. All errors are mine. Email: wooyoung78.lim@gmail.com

1 Introduction

In modern economies, information plays an important role in making appropriate decisions. Unfortunately, the authority to make decisions and the access to decision-relevant information often do not rest with the same party. When an international manufacturing company enters a particular national market, for example, it typically lacks information about the local business conditions, which may be much better understood by domestic companies. When making faculty hiring decisions, a school dean may not have better access to information about the condition of the job market and the qualification of each job candidate than department heads or recruitment committees. In the workplace, workers may have superior information than supervisors about changes in workplace organization, job descriptions, or work flows that would increase firm productivity. Full use of information at the hands of the better-informed party in such situations is socially desirable, but the extent to which they are able to fully utilize the information is constrained by problems of incentive when the interests of the parties involved are not aligned.

In this paper, we consider a model of bargaining over decision rights between an uninformed principal and an informed but self-interested agent.¹ Examples of bargaining over decision rights are abundant within and across the organizational boundaries. When entering a new market, an international manufacturing company may negotiate with a better-informed domestic company over exclusive right to make decisions on pricing, marketing, advertising, distribution, and so on. In a university, a school dean may negotiate with a department head over the right to make a final hiring decision. Workers may negotiate with their supervisors over discretion and responsibility in the workplace. The central question that this paper seeks to answer is, when interests are misaligned between parties, can the bargaining over decision rights lead to an efficient allocation of the rights such that information is fully used in making decisions, and if so how?

In our model, there are two parties, the uninformed principal (a ‘she’) and the informed but self-interested agent (a ‘he’), and one-dimensional decision-making that affects the welfare of both parties under one-dimensional uncertainty.² We restrict our attention to a particular

¹Throughout this paper, the terms decision rights, authority, and controls are used interchangeably.

²This paper adopts the framework used by Crawford and Sobel [16] and Holmström [31, 32] but departs from the standard literature by assuming that the decision itself is noncontractible while decision rights are contractable. It is well known that fully informed decision-making can be achieved when complete contracting is possible. However, as pointed out by Grossman and Hart [25] and Hart and Moore [30], it is typically impossible or too costly to specify all of the feasible contingencies in advance. Empirical support for the transfer of decision rights across firm boundaries has been provided in several recent papers. For example, Lerner and Merges [41] find this phenomenon in pharmaceutical industries, and Arrunada, Garicano, and Vasquez [7] observe the same

mechanism of bargaining over authority: in our benchmark model, we consider the uninformed principal’s optimal price offer for decision rights when the informed agent will simply decide whether to accept or reject the offer. If the price offer is accepted, the agent pays the principal the asking price and makes a decision. Otherwise, the principal retains the decision rights without making any transfer of utility.³ We find that in this benchmark model, it is optimal for the uninformed principal to use the price offer as a device for screening out certain types of agents. In equilibrium, the principal makes a price offer that is accepted by some, but not all, agent types. This means that the principal may end up retaining decision rights while lacking precise information, which leads to an inefficient outcome.

The model is extended to bargaining over decision rights with explicit communication. Specifically, we assume that the informed agent can send a cheap talk message before bargaining begins.⁴ In the benchmark model, we assume that bargaining is *tacit* in the sense that parties can communicate only by making a price offer that directly affects their payoffs. In contrast, and as noted by Crawford [15], real bargaining is usually *explicit*: parties can also communicate by sending nonbinding messages with no direct effects on their payoffs. The main finding is that once communication via cheap talk is allowed, an equilibrium with no inefficient screening by the principal exists even when conflicts of interest are arbitrarily large. The existence of this *ex-post efficient* equilibrium is remarkable because neither tacit bargaining nor communication via cheap talk alone allows both parties to make full use of the agent’s private information in decision-making.

The construction of the ex-post efficient equilibrium includes truth-telling by the agent which eliminates the incentive of the principal to screen out some types of agents. The principal uses a trigger-type strategy that forces the agent to report truthfully: she makes a message-independent price offer that would almost always be accepted by the agent, who

in automobile industries.

³In many negotiations such as one between a school dean and a department head, or one between a supervisor and a worker, it is natural to assume that the party initially holding the decision-making authority also has full bargaining power that allows him/her to make a take-it-or-leave-it offer to the other party. In a companion model of bargaining considered by Lim [42], it is shown that the efficient allocation of decision rights can be supported by one of multiple equilibria when an informed agent has bargaining power, allowing him to make a take-it-or-leave-it price offer. In the current paper, we consider another extreme case in which the principal has bargaining power. Our main finding, along with the finding presented by Lim [42], demonstrates that the efficient allocation of decision rights can be achieved through a set of bargaining protocols regardless of the initial allocation of bargaining power. Moreover, this paper considers a more general class of bargaining protocols that includes a mediator and demonstrates that under some conditions, the bargaining protocols introduced in this paper and by Lim [42] are optimal.

⁴However, our main result is the existence of an ex-post efficient equilibrium that does not depend on the exact timing of the game.

thus fully reveals his private information. By observing the agent’s decision to accept or reject the offer, the principal can determine whether the agent’s report was truthful or not. That is, the principal will consider a rejection to be evidence against full information revelation and will punish the agent by taking an action that affects the agent more negatively than if the agent had told the truth. This *threat of punishment* compels the agent to report true information during the cheap talk stage and accept the equilibrium-path price offer from the principal in the bargaining stage. It turns out that this threat action coincides with the ideal action of a unique type of agent who may reject the equilibrium price offer with positive probability. Consequently, taking the threat action becomes rational for the principal and, more importantly, credible to the agent.

There are several pieces of theoretical evidence demonstrating that such cheap talk messages play an important role in aligning bargainers’ expectations so that they can reach agreement and determine how to share the resulting surplus (Farrell and Gibbons [21] and Matthews [44]). Farrell and Gibbons [21] study a two-stage bargaining game in which talk may be followed by the sealed-bid double auction studied by Chatterjee and Samuelson [12], a well-known model of bargaining under incomplete information. They show that talk can affect the results, in the sense that the equilibrium with cheap talk features bargaining outcomes that would otherwise be impossible in the absence of talk. Matthews [44] considers a specific bargaining situation with a veto threat and shows that an equilibrium exists in which an informed party (proposer) informs the other (chooser) regarding to which of two sets he belongs. This behavior is not an equilibrium behavior in the absence of cheap talk.

We further extend the model in four directions. First, a multidimensional state space is considered. The main result, the existence and characterization of ex-post efficient equilibria, still holds in the multidimensional state space provided the space is compact. Second, we investigate whether such an equilibrium construction can be preserved in the presence of a type-dependent bias and provide the necessary and sufficient conditions for the existence of the ex-post efficient equilibria. Third, we explore the case in which the principal values authority more than the agent does and show that communication can still improve efficiency of bargaining when this asymmetry is sufficiently small. Fourth, we adopt a mechanism design approach and consider a more general bargaining protocol with a neutral mediator. We show that under some conditions, explicit bargaining is an optimal bargaining mechanism that achieves an upper boundary for social welfare in our framework.

This paper contributes to the literature that investigates contractible and transferable decision rights. In contrast to most existing work, this paper allows utilities to be transferable as well in the presence of asymmetric information. Baker, Gibbons, and Murphy [8] and

Aghion, Dewatripont, and Rey [2], Aghion and Tirole [4] consider contractible decision rights in the incomplete contracting framework to study how and when authority can be traded between parties with no asymmetric information. Aghion, Dewatripont, and Rey [3] use the framework of Grossman and Hart [25] to study the transfer of control as a reputation-building device in the presence of asymmetric information. Transferable decision rights in the presence of asymmetric information are extensively studied in the optimal delegation literature (Holmström [31, 32], Alonso and Matouschek [1], Kováč and Mylovanov [35] and Melumad and Shibano [46]), but utilities are not assumed to be transferable.⁵ Dessein [17] compares two coordination mechanisms—cheap talk communication and full delegation—and shows that the uninformed principal prefers not to fully delegate his decision rights to the informed agent when their preferences diverge substantially. However, several papers allow for transferable utilities under asymmetric information. Bester [10] studies a contracting problem in the setting of a monetary transfer when only decision rights are contractible *ex-ante*. Bester [10] focuses on the question of whether a direct and truthful mechanism can implement the same allocation of decision rights as perfect information. Kräbmer [36] considers a message-contingent delegation in which the principal can allocate decision rights after observing cheap talk messages from an informed agent and demonstrates that cheap talk creates incentives for the revelation of information. Krishna and Morgan [40] consider a model in which a principal uses a message-contingent monetary transfer but retains decision rights. These papers consider mechanism design questions, while the current paper studies the equilibria of a bargaining game over decision rights between agent and principal.

This paper is also related to the literature on strategic information transmission.⁶ In their seminal paper, Crawford and Sobel [16] consider a game in which an informed sender sends a nonbinding message that has no direct payoff implication for the receiver, who then takes an action that determines the payoffs of both parties. A clear comparative statics result was found: more information will be transmitted when the preferences between parties are more aligned. Some recent studies have discussed how to facilitate communication between parties and when the transmission of information can be improved. Krishna and Morgan [39], for example, consider the two-stage communication protocol, and Blume, Board and Kawamura [11] consider the possibility of errors in communication. More recently, Goltsman, Hörner, Pavlov and Squintani [24] study communication through a neutral trustworthy mediator and prove that information transmission is improved under the optimal mediation rule. Ivanov [33] demonstrates that an appropriately biased mediator can replace the neutral mediator in Goltsman, Hörner, Pavlov and Squintani [24].

⁵One exception is Ambrus and Egorov [5], who consider money burning in the optimal delegation framework.

⁶For an excellent survey of recent developments in the literature on this subject, see Sobel [52].

The rest of this paper is organized as follows. The next section describes the environment. In Section 3, we establish a benchmark model for bargaining over decision rights. A full characterization of the equilibria is provided, assuming the parties' prior beliefs are *uniform*. In Section 4, we extend the basic model by allowing parties to communicate prior to bargaining by sending cheap talk messages. We demonstrate that an ex-post efficient perfect Bayesian equilibrium exists. We discuss four extensions of the model in Section 5. We conclude the study in Section 6.

2 Model

2.1 Environment

The model contains two parties: a principal (P) and an agent (A). The principal has the initial decision-making authority but little information regarding the state of the world $\theta \in \Theta \equiv [0, 1]$. She has a prior distribution F that is absolutely continuous with density $f > 0$ on its support $[0, 1]$. The agent has different interests from the principal and *privately* knows the true state of the world θ but does not have decision-making authority. The payoffs for a given allocation of authority depend on the action $y \in Y$ taken by the party with decision-making authority and the state of the world θ . The payoff functions of the parties take the forms $U^P(y, \theta) = -l(|y - \theta|)$ for the principal and $U^A(y, \theta, b) = -l(|y - (\theta + b)|)$ for the agent.⁷ We refer to l as the loss function and assume that $l''(\cdot) > 0$, $l'(0) = 0$ and $l(0) = 0$. In this model, the ideal action of the principal is $\bar{y}^P(\theta) = \theta$, and the ideal action of the agent is $\bar{y}^A(\theta, b) = \theta + b$, where $b > 0$ is a parameter that measures how nearly the agent's interest coincides with that of the principal. Assume $\{\theta, \theta + b\} \in Y$ for any $\theta \in \Theta$. All of these conditions are common knowledge between the parties.

2.2 Ex-post efficiency

Define an *ex-post* efficient action as follows. An action is said to be *efficient ex-post* if there is no other feasible action that makes some individual better off without making other individuals worse off after the true state of the world θ is publicly known.

Definition 1. *An action $y \in \mathbb{R}$ is efficient ex-post at θ if there is no other action $z \in \mathbb{R}$ such that*

$$U^P(z, \theta) \geq U^P(y, \theta) \quad \text{and} \quad U^A(z, \theta, b) \geq U^A(y, \theta, b) \quad (1)$$

⁷A special case is quadratic utilities ($U^P(y, \theta) = -(y - \theta)^2$ and $U^A(y, \theta, b) = -(y - \theta - b)^2$), which we assume in most examples and applications. Similar utility functions are assumed in many other studies: for example, Dessein [17]. To see more papers assuming quadratic utilities, see Kováč and Mylovánov [35].

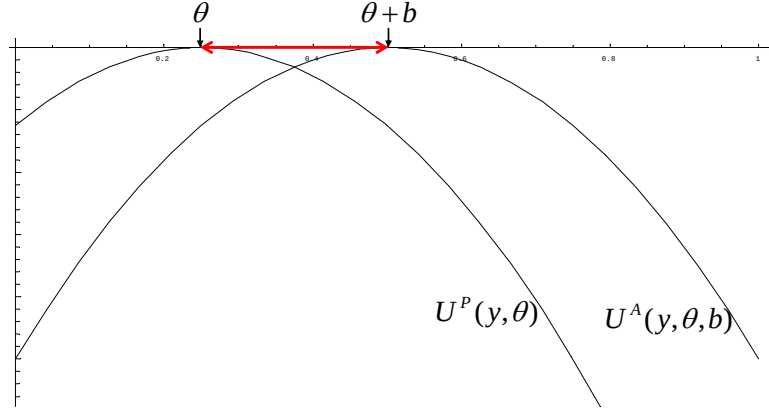


Figure 1: Ex-post Efficient Actions

with at least one strict inequality.

In our environment, an action y is efficient *ex-post* if and only if $y \in [\theta, \theta + b]$ when the realization of the state is θ , as one can see in Figure 1. Based on the notion of *ex-post* efficient action, we define the ex-post efficiency as follows.

Definition 2. An outcome $O : \Theta \rightarrow Y$ is **ex-post efficient** if $O(\theta)$ is efficient *ex-post* for any $\theta \in \Theta$.

It is worth noting that this notion of ex-post efficiency is different from the standard notion of efficiency in the environment with quasi-linear utilities.⁸ The notion of ex-post efficiency captures whether the allocation of decision-rights is efficient or not in our environment where decision is not contractible.

3 Benchmark: Tacit Bargaining

Consider bargaining between the informed agent and the uninformed principal over decision-making authority. The timing of the game is as follows:

1. The agent privately observes the state of the world $\theta \in \Theta \equiv [0, 1]$.
2. The principal makes an offer $p \in \mathbb{R}$ to transfer the authority to take an action.⁹
3. The agent decides whether to accept or reject the offer.
4. If the agent accepts the offer, he pays the price to the principal and takes an action,

⁸According to the standard efficiency in the environment with quasi-linear utilities, an action is efficient *ex-post* at θ if it maximizes the sum of utilities of the two parties.

⁹We allow p to be negative, which means that the principal pays $|p|$ to the agent.

denoted by y^A . In this case, the payoffs of the principal and the agent are $U^P(y^A, \theta) + p$ and $U^A(y^A, \theta, b) - p$, respectively. If the agent rejects the offer, however, the principal takes an action y^P without transferring the decision-making authority. The payoffs of the principal and the agent are then $U^P(y^P, \theta)$ and $U^A(y^P, \theta, b)$, respectively.

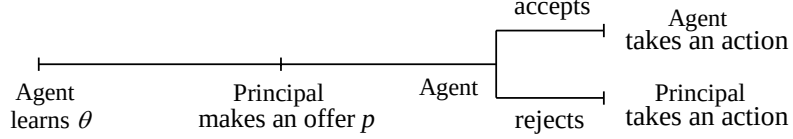


Figure 2: Timing of the game

We use the concept of a perfect Bayesian equilibrium to define our equilibrium. The principal's strategy consists of a price offer p^* and an action rule y^P . The action rule, denoted by $y^P : \mathbb{R} \rightarrow \mathbb{R}$, specifies the principal's action after the rejection of each price offer $p \in \mathbb{R}$ that she might make. Because the utility function is strictly concave in y , the principal will never use mixed strategies in equilibrium. The agent's strategy consists of a decision rule and an action rule. The decision rule, denoted by $d^A : \Theta \times \mathbb{R} \rightarrow [0, 1]$, specifies a probability of rejection for each price offer $p \in \mathbb{R}$ that he might receive. The action rule, $y^A : \Theta \times \mathbb{R} \rightarrow \mathbb{R}$, specifies the agent's choice of action after he accepts the principal's price offer p . Without loss of generality, we can focus on a pure strategy action rule due to the strict concavity of the utility function. The strategy profile $\{(p^*, y^P), (d^A, y^A)\}$ and the principal's belief $\rho(\theta|p)$ form a perfect Bayesian equilibrium if

(B1) p^* solves

$$\max_{p \in \mathbb{R}} \int_0^1 \{d^A(\theta, p)U^P(y^P(p), \theta) + (1 - d^A(\theta, p))(p - l(b))\} f(\theta) d\theta$$

(B2) for each $p \in \mathbb{R}$ and each $\theta \in [0, 1]$, $d^A(\theta, p)$ solves

$$\max_{d^A \in [0, 1]} (1 - d^A)(-p) + d^A \cdot U^A(y^P(p), \theta, b)$$

(B3) for each $\theta \in [0, 1]$ and $p \in \mathbb{R}$, $y^A(\theta, p) = \bar{y}^A(\theta) = \theta + b$

(B4) for each $p \in \mathbb{R}$, $y^P(p)$ solves

$$\max_{y \in \mathbb{R}} \int_0^1 U^P(y, \theta) \rho(\theta|p) d\theta$$

where $\rho(\theta|p)$ is the principal's updated belief after observing the agent's rejection of p , which is provided by Bayes' rule whenever possible.

3.1 Equilibrium

3.1.1 Example: *uniform* distribution

In this section, we illustrate the primary idea behind our analysis assuming that f is *uniform* over $[0, 1]$. This assumption, along with quadratic utility functions, follows a leading example presented by Crawford and Sobel [16] that has been widely used in the literature on strategic information transmission and optimal delegation. We will extend our result to more general distributions in the next subsection. In the following sections, we first focus on the agent's decision rule, which satisfies a condition called *monotonicity*. Next, with a full characterization of the agent's decision rule, we show there exists a unique price offer that maximizes the principal's expected utility.

For an arbitrary $p \in \mathbb{R}$, define the set of agent types who accept p with probability one as

$$\Theta(p) = \{\theta \in [0, 1] | d(\theta, p) = 0\}.$$

Define the set of agent types who reject the offer p with probability one as

$$\Theta'(p) = \{\theta \in [0, 1] | d(\theta, p) = 1\}.$$

An agent's decision rule is said to be *monotonic* if for any price offer for which there is an agent type θ who accepts the offer with positive probability, all agent types higher than θ accept it with probability one. This monotonicity condition is assumed to characterize the set of equilibria in this section, but it will be formally proven to be necessary in the next subsection for more general distributions.

This monotonicity implies that for any $p \in \mathbb{R}$, both $\Theta(p)$ and $\Theta'(p)$ are convex and ensures that for any $p \in \mathbb{R}$, there is at most one agent type that is indifferent between acceptance or rejection of the offer. Let $\theta_p \in [0, 1]$ denote this agent type if it exists. Then we can state that $\Theta(p) = (\theta_p, 1]$ and $\Theta'(p) = [0, \theta_p)$. From the indifference condition at θ_p , we have

$$p = l(|y^P(p) - \theta_p - b|), \tag{2}$$

where $y^P(p) = \arg \max_y \int_0^{\theta_p} -l(|y - \theta|) \cdot \frac{1}{\theta_p} d\theta = \frac{\theta_p}{2}$. Thus, we have

$$p = l(\frac{\theta_p}{2} + b) \quad \text{or} \quad \theta_p = 2(l^{-1}(p) - b). \tag{3}$$

Because $\theta_p \in [0, 1]$, we have the following corollary:

Corollary 1.

$$\Theta(p) = \begin{cases} [0, 1] & \text{if } p < l(b), \\ (2(l^{-1}(p) - b), 1] & \text{if } l(b) \leq p \leq l(b + \frac{1}{2}), \\ \emptyset & \text{if } p > l(b + \frac{1}{2}) \end{cases}$$

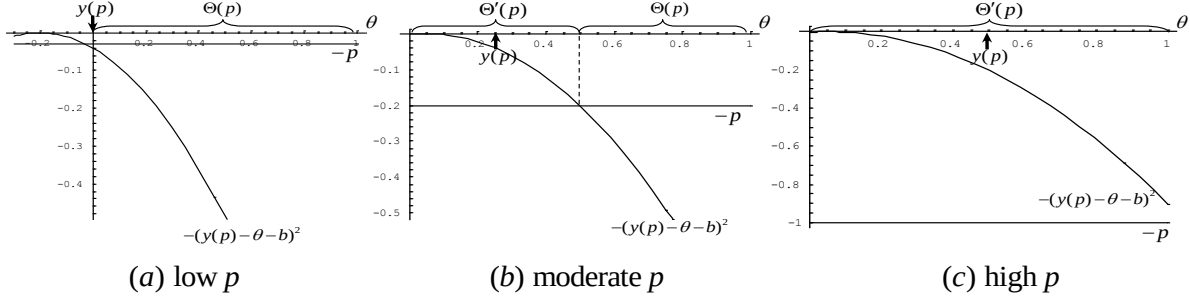


Figure 3: The agent's decision rule

In words, all agent types in $[0, 1]$ accept a low price offer ($p < l(b)$) with probability one, and once the price offer becomes greater than $l(b)$, these low agent types start rejecting it. As p increases, the set $\Theta(p)$ becomes smaller, until finally all of the agent types in $[0, 1]$ reject a high price offer ($p > l(b + \frac{1}{2})$) with probability one.

In Figure 3, we see that the principal is faced with a clear trade-off between a higher price offer and a higher chance of rejections. Although a higher price offer provides a higher payoff to the principal if accepted, it does not appear to be optimal for the principal to make such an offer because it is unlikely to be accepted (Figure 3(c)). Similarly, making a price offer that could be accepted by all agent types does not appear to be optimal because the offer must be sufficiently low (Figure 3(a)). These observations suggest that the principal's optimal price offer should lie halfway between these two extremes (Figure 3(b)). To confirm this idea, consider the principal's optimal price offer as a best response to the agent's strategy.¹⁰ The principal chooses p^* to solve

$$\begin{aligned} \max_{p \in \mathbb{R}} EU^P &= \int_0^{\theta_p} -l(|y^P(p) - \theta|) d\theta + (1 - \theta_p)(p - l(b)) \\ \text{s.t. } \theta_p &= \begin{cases} 0 & \text{if } p < l(b), \\ 2(l^{-1}(p) - b) & \text{if } l(b) \leq p \leq l(b + \frac{1}{2}), \text{ and } y^P(p) = \frac{\theta_p}{2}. \\ 1 & \text{if } p > l(b + \frac{1}{2}) \end{cases} \end{aligned} \quad (4)$$

Note that there is a unique interior solution to this maximization problem due to the strict concavity of EU^P in p . From the first-order condition, we obtain

$$p^* = l(\frac{1}{4} + b) \quad \text{and} \quad \theta_{p^*} = \frac{1}{2}. \quad (5)$$

This implies that it is optimal for the principal to make a price offer acceptable to some agents

¹⁰When p is so low that no agent types reject it, the principal's choice of action $y^P(p)$ is off the equilibrium path. In this case, we take $y^P(p) = 0$ to support our equilibrium construction.

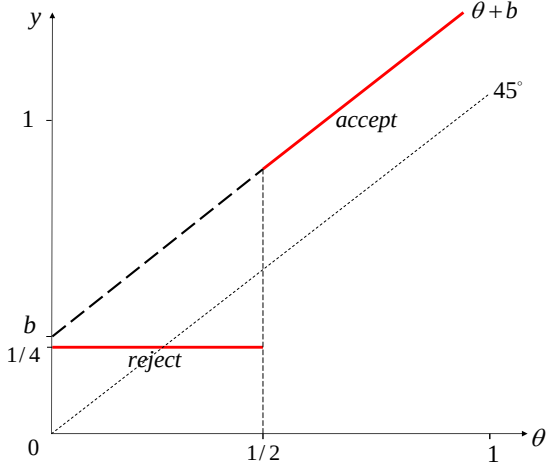


Figure 4: Equilibrium Outcome

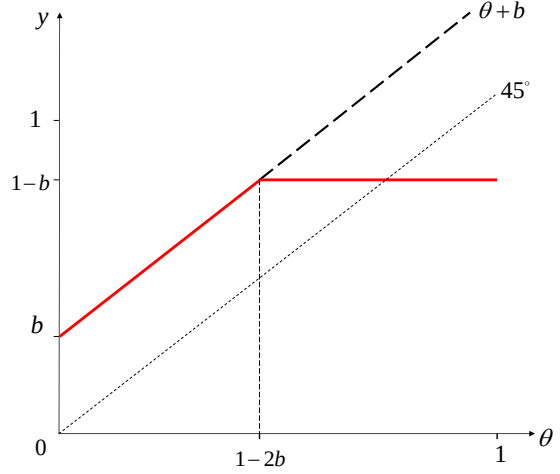


Figure 5: Optimal Delegation

of the high type but not to the remaining agents of the low type. This result is summarized in the following proposition.

Proposition 1. *Under tacit bargaining with a uniform prior, there is a unique equilibrium outcome in which the principal makes a price offer $p^* = l(\frac{1}{4}+b)$, the agent type $\theta \in [0, \frac{1}{2})$ rejects the offer with probability one, and the agent type $\theta \in (\frac{1}{2}, 1]$ accepts the offer with probability one.*

This outcome may appear to be the opposite of the outcome of optimal delegation as studied by Holmström [31][32], Melumad and Shibano [46], Alonso and Matouschek [1], Goltsman *et al.* [24] and Kováč and Mylovánov [35]. According to the optimal delegation rule in the uniform-quadratic environment, the informed agent can enforce any decision he likes as long as it does not exceed $1 - b$ (see Figure 5). It is beneficial for the principal to impose an upper bound on the allowable actions because the informed party's ideal action is always higher than that of the principal. In our equilibrium, by contrast, the agent can enforce any decision he likes as long as it exceeds $\frac{1}{2} + b$ (see Figure 4). The following justification can be provided for this result. Obviously, the principal can maximize her *ex-ante* utility by setting her price based on the agent types that exhibit higher willingness-to-pay for decision rights. An agent's willingness-to-pay is measured by the distance between his ideal action and the action that the principal would take if the agent rejects the price offer. Because the informed agent's ideal action is always higher than that of the principal given any action taken by the principal after a rejection, agent whose types are higher than the action taken by the principal are willing to pay more to obtain decision rights than agent whose types are lower.

3.1.2 General Case

We now extend the analysis of the previous subsection to more general distributions. Recall that in the previous section, the monotonicity of the agent's decision rule allows us to obtain a unique optimal price offer for the principal. The following regularity condition on the parties' prior belief f is sufficient to enforce the same monotonicity in the agent's decision rule and, as a result, ensures that all of the qualitative results obtained in the previous section are preserved.

Condition 1. *For any given value of $b > 0$,*

$$y(\underline{\theta}, \bar{\theta}) - b < \frac{\underline{\theta} + \bar{\theta}}{2} \quad (6)$$

for any $\underline{\theta}$ and $\bar{\theta}$ with $0 \leq \underline{\theta} \leq \bar{\theta} \leq 1$, where

$$y(\underline{\theta}, \bar{\theta}) = \begin{cases} \operatorname{argmax}_{\theta} \int_{\underline{\theta}}^{\bar{\theta}} U^P(y, \theta) f(\theta) d\theta & \text{if } \underline{\theta} < \bar{\theta}, \\ \bar{\theta} & \text{if } \underline{\theta} = \bar{\theta}. \end{cases}$$

In other words, this condition implies that for any interval subset of Θ , the ideal action of a principal who believes that agent types within the interval will reject a price offer is not excessively weighted towards the right of the interval.¹¹ Any prior f satisfies this regularity condition if $b \geq 1/2$. Moreover, this condition holds for any $b > 0$ if f is nonincreasing in θ . In particular, this condition is satisfied given the *uniform* distribution considered in the previous subsection.

Under this condition of regularity, the agent's decision rule satisfies the *monotonicity* condition.

Lemma 1 (Monotonicity). *If Condition 1 is satisfied, then the agent must use the monotonic decision rule. That is, for any price offer, if there exists an agent type θ who accepts the offer with positive probability, then all agent types higher than θ must accept it with probability one.*

Proof. See Appendix B. □

As we have already discussed in the case of the uniform distribution, monotonicity makes it optimal for the principal to use her price offer as a screening device. This implies that the principal sometimes retains decision rights while lacking precise information, which leads to an inefficient outcome. This result is summarized in the following proposition.

¹¹For example, take b sufficiently small and suppose that the ideal action of the principal after rejection is sufficiently close to the indifference agent type θ_p . From the indifference condition of θ_p and the symmetry of the loss function $l(\cdot)$, there may be an agent type to the left of θ_p that strictly prefers to accept the offer. Thus, the monotonicity condition may not hold.

Proposition 2. *Under tacit bargaining, any equilibrium where the agent uses a monotonic decision rule leads to an ex-post inefficient outcome.*

4 Explicit Bargaining

In this section, we explore how introducing explicit communication into the basic model affects the outcomes. The timing of the game is as follows:

1. The agent privately observes the state of the world $\theta \in \Theta \equiv [0, 1]$.
2. The agent sends a message $m \in M \equiv [0, 1]$ to the principal.¹²
3. After observing the message from the agent, the principal makes a price offer $p \in \mathbb{R}$ to transfer authority to take an action.
4. The agent decides whether to accept or reject the offer.
5. If the agent accepts the offer, he pays the price offered by the principal and takes an action, denoted by y^A . In this case, the payoffs of the principal and the agent are $U^P(y^A, \theta) + p$ and $U^A(y^A, \theta, b) - p$, respectively. If the agent rejects the offer, however, the principal takes an action denoted by y^P without transferring the decision-making authority. The payoffs of the principal and the agent are then $U^P(y^P, \theta)$ and $U^A(y^P, \theta, b)$, respectively.

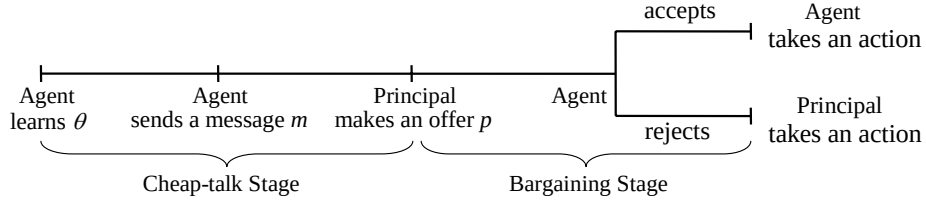


Figure 6: Communication Before Bargaining

Again, we use the concept of a perfect Bayesian Equilibrium to define our equilibrium. The agent's strategy consists of a message rule, a decision rule and an action rule. The message rule $\mu : \Theta \rightarrow \Delta(M)$ specifies the choice of message for each type $\theta \in \Theta$. The decision rule, denoted by $d : \Theta \times M \times \mathbb{R} \rightarrow [0, 1]$, specifies the probability of rejection for each price offer $p \in \mathbb{R}$ that the agent who sent the message m might receive. The action rule, $y^A : \Theta \times M \times \mathbb{R} \rightarrow \mathbb{R}$, specifies the action taken by the agent type θ who sent the message m and accepted the principal's price offer p . The principal's strategy consists of a price rule and an action rule. The price rule $p^* : M \rightarrow \mathbb{R}$ specifies the principal's choice of price offer for each message $m \in M$.

¹²For simplicity, we assume $M = [0, 1]$. We could derive the same result by assuming that M is any Borel measurable set that has a cardinality at least as large as the type space.

that the principal might receive. The action rule, denoted by $y^P : M \times \mathbb{R} \rightarrow \mathbb{R}$, specifies the action taken by the principal who observed a message m and submitted a price offer p that was rejected. The strategy profile $\{(\mu, d, y^A), (p^*, y^P)\}$ and the principal's posterior beliefs ρ_1 and ρ_2 form a perfect Bayesian equilibrium if

(CB1) for each $\theta \in [0, 1]$, $\int_M \mu(m|\theta)dm = 1$, and if $m^* \in M$ is in the support of $\mu(\cdot|\theta)$, then m^* solves

$$\max_{m \in M} d(\theta, m, p^*(m))U^A(y^P(m, p^*(m)), \theta, b) - p^*(m)(1 - d(\theta, m, p^*(m)))$$

(CB2) for each $m \in M$, $p^*(m)$ solves

$$\max_{p \in \mathbb{R}} \int_0^1 \{d(\theta, m, p)U^P(y^P(m, p), \theta) + (1 - d(\theta, m, p))(p - l(b))\}\rho_1(\theta|m)d\theta$$

(CB3) for each $\theta \in [0, 1]$, $m \in M$, and $p \in \mathbb{R}$, $d(\theta, m, p)$ solves

$$\max_{d \in [0, 1]} (1 - d)(-p) + d \cdot U^A(y^P(m, p), \theta, b)$$

(CB4) for each $\theta \in [0, 1]$, $m \in M$, and $p \in \mathbb{R}$, $y^A(\theta, m, p) = \theta + b$

(CB5) for each $m \in M$ and $p \in \mathbb{R}$, $y^P(m, p)$ solves

$$\max_{y \in \mathbb{R}} \int_0^1 U^P(y, \theta)\rho_2(\theta|m, p)d\theta$$

(CB6)

$$\rho_1(\theta|m) = \frac{\mu(m|\theta)}{\int_0^1 \mu(m|\theta')d\theta'} \quad \text{and} \quad \rho_2(\theta|m, p) = \frac{d(\theta, m, p)\rho_1(\theta|m)}{\int_0^1 d(\theta', m, p)\rho_1(\theta'|m)d\theta'}$$

where $\rho_1(\theta|m)$ is the principal's updated belief after observing the message m from the agent and $\rho_2(\theta|m, p)$ is the principal's updated belief after receiving the message m and observing the rejection of the price offer p . These updated beliefs are provided by Bayes' rule whenever possible.

4.1 Ex-post Efficient Equilibria

Is it possible that communication before bargaining is informative and thus improves the efficiency of bargaining? In this section, we will show that there exists a perfect Bayesian equilibrium in which pre-bargaining communication is used to achieve the ex-post efficient outcome by preventing the principal from inefficiently screening the agent types.

Consider the following strategy profile: the principal makes a price offer $l(b)$ regardless of the message he observed and takes an action $y = \bar{y}^P(0) = 0$ if the offer is rejected. The

agent fully reveals his private information by sending a truth-telling message $m = \theta$ during the cheap talk stage and accepts any offer less than $l(b)$ with probability one but rejects any offer strictly greater than $l(b)$ with probability one. Agent type $\theta \in (0, 1]$ accepts price offer $l(b)$ with probability one, while agent type $\theta = 0$ randomizes between accepting and rejecting the offer. If the principal makes any price offer $p \neq l(b)$ and is rejected, the principal takes an action $y = m$ following the message from the agent in the cheap talk stage. It is easy to see that the agent types' strategies represent the best responses to the principal's strategy. First, no agent type has an incentive to deviate in the cheap talk stage because the principal's price offer is message-independent. Second, no agent type has an incentive to reject the offer in the bargaining stage because for all $\theta \in [0, 1]$

$$\underbrace{-l(b)}_{\text{from accepting } l(b)} \geq \underbrace{-l(|0 - \theta - b|)}_{\text{from rejecting } l(b)} = -l(\theta + b).$$

Given the agent's strategy specified above, the principal's best response is to make an offer $l(b)$, the highest price offer accepted by agent types who tell the truth in the cheap talk stage. The principal does not have an incentive to make any offer $p < l(b)$ because such price offers will be accepted by all agent types and provide a strictly lower payoff to the principal than making the price offer $l(b)$. The principal does not have a strict incentive to make any offer $p > l(b)$ because such price offers will be rejected by all agent types, and the actions that would be taken by the principal after $p > l(b)$ is rejected will provide the same payoff as could be obtained by making the offer $p = l(b)$. After the price offer $l(b)$ is rejected, the principal believes that the true state of the world is $\theta = 0$ with probability one. This is a reasonable belief, in the sense that the agent type $\theta = 0$ is the only type who is indifferent between accepting and rejecting the offer $l(b)$ and all other agent types strictly prefer accepting the offer. We will discuss this issue more carefully in Section 6.

Proposition 3 (Ex-post Efficient Equilibrium). *For any $b > 0$, there exists an ex-post efficient perfect Bayesian equilibrium. In this equilibrium, the informed agent fully reveals his private information by sending truth-telling messages in the cheap talk stage, and the principal makes a message-independent price offer that is accepted by (almost) all of the agent types, resulting in fully informed decisions being made by the agent.*

Proof. See Appendix B. □

Note that in this equilibrium, the information revealed is not used directly by the principal, but it does determine the agent's option value of rejecting a price offer, which is $-l(b)$. Given the full revelation of information, the option value of rejecting a price offer becomes

independent of agent type, which allows the principal to find a single price offer that will be (almost) always accepted. This situation would have been impossible without full disclosure of information in the cheap talk stage. Under partial disclosure, different agent types should have had different option values for rejecting a price offer. The values would be determined by the actions that would be taken by the principal after rejection of the price offer.

In this equilibrium, decision rights are allocated efficiently, so that the information available to the agent is fully utilized. As a result, the equilibrium outcome is efficient ex-post. These efficient allocations of decision rights cannot be attained without communication, as our previous analysis suggested for tacit bargaining. Without communication, the principal would want to use the price offer to inefficiently screen out some agent types.

It is worth noting that the existence of the ex-post efficient equilibrium is robust against the exact timing of the game. To be more precise, consider the game in which bargaining proceeds first and communication follows under the contingency that an agreement is not reached. The following perfect Bayesian equilibrium then exists in this game and is outcome-equivalent to the ex-post efficient equilibrium in the original model. The principal makes a price offer $l(b)$ and takes an action $y = 0$, regardless of the message she receives from the agent if the offer is rejected. All agent types in Θ but $\theta = 0$ accept the offer $l(b)$ with probability one and fully reveal their private information by sending truth-telling messages off the equilibrium path (i.e., when any offer is rejected.) Because it is straightforward to conclude that this strategy profile satisfies the mutual best responses under certain beliefs derived by Bayes' rule, I omit the detailed proof here.

5 Discussion and Extensions

This section is devoted to a discussion on the robustness of the ex-post efficient equilibria to four possible extensions of the model. First, a multidimensional state space is considered. We show that as long as the state space is compact, the dimensionality does not matter in the construction of the ex-post efficient equilibria. Second, we examine the case of type-dependent bias and provide the necessary and sufficient conditions for the existence of the ex-post efficient equilibria. Third, we explore the case in which the principal values authority more than the agent does, and show that communication can still improve efficiency of bargaining when this asymmetry between the principal and the agent is sufficiently small. Fourth, we adopt a mechanism design approach and show that under some conditions, explicit bargaining is an optimal bargaining mechanism.

In Appendix A, we also discuss equilibrium refinement and robustness of the ex-post efficient equilibria. We apply two standard cheap talk refinements, *neologism-proofness* (Farrell

[18]) and NITS (Chen, Kartik and Sobel [13]), and show that the existence of the ex-post efficient equilibria is robust against those refinements: no matter how much two parties' preferences differ, the equilibrium is neologism-proof in Farrell's [18] sense. Imposing NITS (no incentive to separate), the criterion proposed by Chen, Kartik and Sobel [13] to refine equilibria in cheap talk games (Crawford and Sobel [16]) leads to the same result. We also discuss the notions of sequential equilibrium proposed by Kreps and Wilson [38] and extensive-form trembling-hand perfection described by Selten [51]. Interested readers are referred to Appendix A.

5.1 Multidimensional State Space

In this section, we extend our model to a multidimensional state space. We show that the compactness of the state space plays a crucial role in the existence of the ex-post efficient equilibrium, but the dimensionality does not. Suppose that the state space Θ is a compact subset of $\mathbb{R}^d$. Similarly, the action space Y is a subset of $\mathbb{R}^d$. $b \in \mathbb{R}^d$ is the bias of the agent. Assume that for any $\theta \in \Theta$, $\{\theta, \theta + b\} \subseteq Y$. For state θ and action y , the principal's utility is $-l(|y - \theta|)$, and the agent's utility is $-l(|y - \theta - b|)$, where $|\cdot|$ is the Euclidean norm.

It is obvious that the construction of the ex-post efficient equilibrium described in the previous section can be extended to the current model with the multidimensional state space. Let $b \cdot \theta$ denote the inner product of b and θ . Because $b \cdot \theta$ is continuous and Θ is compact, the minimizer of $b \cdot \theta$, which will be denoted by z , exists in Θ . The following strategy profile constitutes a perfect Bayesian equilibrium: (i) All agent types fully reveal their types in the cheap talk stage. (ii) All agent types except z accept any offer $p \leq l(|b|)$ with probability one; otherwise, they reject it with probability one. (iii) The agent type z rejects the offer $p = l(|b|)$ with positive probability and behaves in the same way as other agent types for any other price offers. (iv) The principal makes a price offer $l(|b|)$ regardless of the message received from the agent. (v) If the offer $p = l(|b|)$ is rejected, the principal believes that the agent's type is z with probability one and takes an action $y = \bar{y}^P(z) = z$. (vi) If the offer $p \neq l(|b|)$ is rejected, the principal believes the messages from the agent and takes the action $y = m$. (vii) If any offer is accepted, the agent type θ takes his ideal action $\theta + b$.

This equilibrium is graphically illustrated in Figure 7. In this figure, the solid circle represents the indifference curve of the agent at the threat action z that will be taken by the principal if the equilibrium price offer $l(|b|)$ is rejected. The distance between θ and $\theta + b$ represents the equilibrium price offer $l(|b|)$. As the figure shows, any agent type $\theta \neq z$ must accept the equilibrium price offer because the distance between θ and $\theta + b$ is shorter than the radius of the solid circle with the origin at $\theta + b$. Remarkably, those two coincide

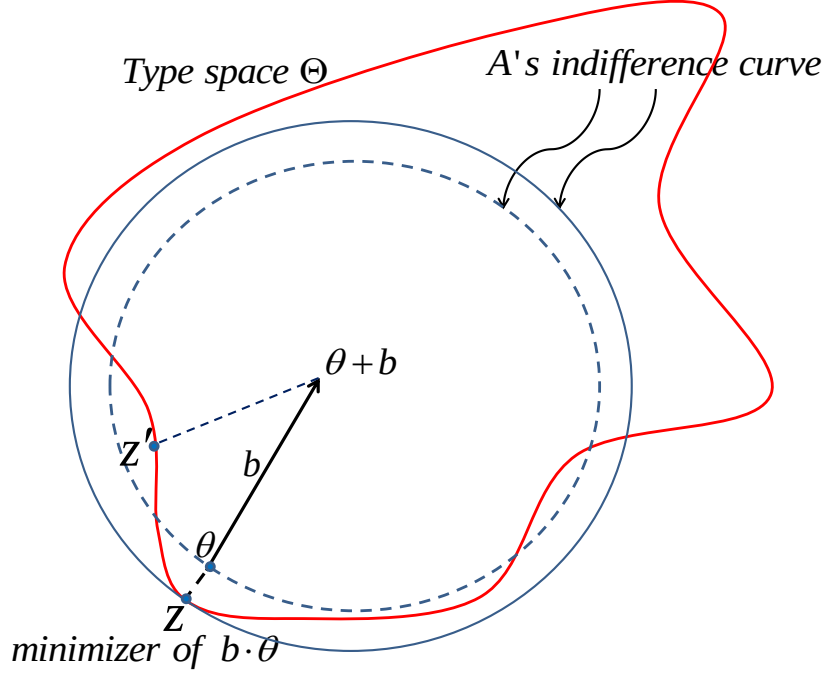


Figure 7: Constructing an ex-post efficient equilibrium

only if $\theta = z$, the minimizer of $b \cdot \theta$, which means that agent type z is indifferent between accepting and rejecting $l(|b|)$. Taking the action z turns out to be rational for the principal (and as a result, credible to the agent) under the belief that the agent type z is the only type who randomizes between accepting and rejecting the price offer. Again, truth telling is optimal for the agent because the principal makes a message-independent price offer $l(|b|)$. The dotted circle represents the indifference curve of the agent at the action $y = \theta$. Truth telling guarantees the agent a utility equal to that of the dotted indifference curve, so a truthful agent will always reject a price offer higher than $l(|b|)$. This makes it optimal for the principal to make such an offer $l(|b|)$ regardless of the messages received.

Figure 7 also emphasizes that the construction of the credible threat action relies heavily on the compactness of the state space. If we take any $z' \neq z$ as a threat action, we can always find an agent type $\theta \in \Theta$ (sufficiently close to z) that strictly prefers z' over truth-telling. The minimizer of $b \cdot \theta$ exists if the type space is compact. This result is summarized in the following proposition.

Proposition 4. *Suppose that $\Theta \subset \mathbb{R}^d$ is compact. For any $b \in \mathbb{R}^d$, there exists an ex-post efficient equilibrium.*

Proof. See Appendix B. □

It is worth noting that full revelation in this game occurs with a single sender, which is not possible in a standard cheap talk game independent of the dimensionality of the type space (Battaglini [9] and Ambrus and Takahashi [6]).

5.2 Type-Dependent Bias

One can easily see that the assumption of constant bias (i.e., the agent's bias does not depend on his type) plays an important role. The construction of the ex-post efficient equilibrium relies on the message-independent (and thus type-independent) price offer. The price is such that it exactly compensates for the loss incurred when the other party makes its preferred decision. Therefore, it is necessary to discuss to what extent this result can be generalized when the bias depends on the agent's type. In this section, assume that $b(\cdot) : \Theta \rightarrow R$ and $|b(\cdot)|$ is *continuous*. All other elements remain the same.

Consider an ex-post efficient equilibrium candidate in which the principal makes a price offer $p = l(|b(0)|)$ on the equilibrium path; whenever observing rejection of her offer, she takes action $y = 0$. The agent type $\theta = 0$ is then indifferent between accepting and rejecting this offer. For agent type $\theta > 0$, the payoff from accepting $p = l(|b(0)|)$ is

$$-l(|b(0)|) - l(|\theta + b(\theta) - \theta - b(\theta)|) = -l(|b(0)|)$$

and the payoff from rejecting it is

$$-l(|0 - \theta - b(\theta)|) = -l(|\theta + b(\theta)|)$$

Therefore, as long as the following is true, any agent type $\theta \in (0, 1]$ has incentives to accept the equilibrium path price offer:

$$-l(|b(0)|) \geq -l(|\theta + b(\theta)|), \quad \forall \theta \in [0, 1]. \quad (7)$$

Observe that the principal has already received a truthful message when she makes a price offer on the equilibrium path. Therefore, she is willing to sell decision rights at price $l(|b(0)|)$ only if

$$\underbrace{0}_{\text{from taking her ideal action herself}} \leq \underbrace{l(|b(0)|) - l(|b(\theta)|)}_{\text{from selling the decision rights to the agent}}, \quad \forall \theta \in [0, 1]$$

or, equivalently,

$$-l(|b(0)|) \leq -l(|b(\theta)|), \quad \forall \theta \in [0, 1]. \quad (8)$$

From (7), (8), and $l(|\cdot|)' \geq 0$, we obtain the following condition:

$$|\theta + b(\theta)| \geq |b(0)| \geq |b(\theta)|, \quad \forall \theta \in [0, 1]. \quad (9)$$

The second weak inequality implies that it is necessary for 0 to be a global maximizer of $|b(\cdot)|$. The first weak inequality states that the gap between $|b(0)|$ and $|b(\theta)|$ must be bounded by θ .¹³ For example, these conditions will be satisfied with $b(\theta) = 1 - \frac{\theta}{2}$ when using a quadratic loss function because

$$-(1 + \frac{\theta}{2})^2 \leq -1 \leq -(1 - \frac{\theta}{2})^2, \quad \forall \theta \in [0, 1].$$

Given the condition (9), the principal does not wish to deviate to a higher price because it would be rejected by all agent types. This occurs because for every off-the-equilibrium-path price offer, (especially $p > l(|b(0)|)$), if the offer is rejected, the principal takes an action $y = \theta$ but $y = 0$. This action choice $y = \theta$ is optimal given her belief. Each agent type θ will then compare the payoff from rejection, $-l(|\theta - \theta - b(\theta)|) = -l(|b(\theta)|)$, and the payoff from acceptance, $-p - l(|\theta - b(\theta) - \theta - b(\theta)|) = -p < -l(|b(0)|)$. Under the condition (9), the payoff from rejection is (weakly) higher than the payoff from acceptance for all $\theta \in [0, 1]$.

More generally, it is clear that for a global maximizer of $|b(\cdot)|$, denoted by $\bar{\theta} \in (0, 1)$ ¹⁴, we can construct an ex-post efficient equilibrium with a (credible) threat action $y = \bar{\theta}$ and an equilibrium path price offer $p = l(|b(\bar{\theta})|)$ if

$$|\bar{\theta} - \theta - b(\theta)| \geq |b(\bar{\theta})|, \quad \forall \theta \in [0, 1]. \quad (10)$$

This condition, along with a simple triangle inequality, implies that the gap between $|b(\bar{\theta})|$ and $|b(\theta)|$ must be bounded by $|\theta - \bar{\theta}|$, which requires discontinuity in $b(\cdot)$ at $\bar{\theta}$. Figure 8 demonstrates an example of $b(\cdot)$ that satisfies the condition (10). The results are summarized in the following proposition.

Proposition 5. *For a continuous $b(\cdot)$, there exists an ex-post efficient equilibrium with a threat action $z \in \{0, 1\}$ if and only if*

$$|\theta + b(\theta)| \geq |b(z)| \geq |b(\theta)|, \quad \forall \theta \in [0, 1].$$

Proof. See Appendix B. □

¹³It is obvious that the necessary and sufficient condition for action $y = 1$ to be a credible threat action is $|1 - \theta - b(\theta)| \geq |b(1)| \geq |b(\theta)|$, for any $\theta \in [0, 1]$.

¹⁴Because $|b(\cdot)|$ is continuous and Θ is compact, the global maximizer of $|b(\cdot)|$ exists in Θ .

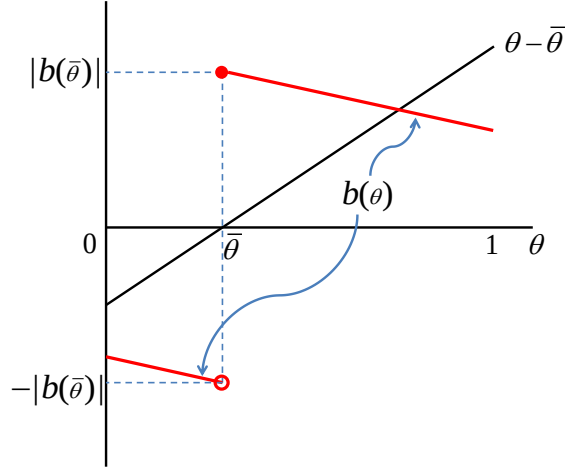


Figure 8: Type-dependent Bias

5.3 Value of Authority and Value of Communication

In this section, we consider the situation in which the principal values the authority to make a decision more than the agent does. In organizations, decisions often matter more for principals than for agents. That is, the principal incurs larger losses than the agent from equivalent deviations from their most preferred action. In this case, our analysis in the previous sections suggests that the ex-post efficient equilibria may not exist.

For expositional simplicity, we assume that the prior is uniform and the payoff functions are

$$U^P(y, \theta) = -\alpha \cdot (y - \theta)^2 \quad \text{and} \quad U^A(y, \theta, b) = -(y - \theta - b)^2 \quad (11)$$

for the principal and the agent respectively, where $\alpha > 1$ is a parameter that measures how much more important the decision is to the principal than to the agent.¹⁵ We would like to address mainly the following two questions: 1) Is it still possible to have an equilibrium in which fully informed decisions are made with probability one? 2) Otherwise, how does the higher value of authority for the principal affect the equilibrium outcome?

Observe that the construction of the ex-post efficient equilibria in the previous sections does not work when $\alpha > 1$ because the fully informed principal would not be willing to give up the decision rights at the price $l(b)$. In general, full revelation cannot be supported in equilibrium because the fully informed principal would make an unacceptable offer and exercise the decision right, which creates incentives for the agent to modify information in the cheap talk stage. Without full revelation, full delegation cannot occur either because under

¹⁵It is straightforward to see that when $\alpha < 1$ the ex-post efficient equilibria survive.

partial disclosure different type of agents would have had different values for rejecting a price offer as we discussed in Section 4.1 such that the principal cannot find a single price offer that will (almost) always be accepted.

When $\alpha > 1$ is sufficiently small, however, there exists an equilibrium in which communication can improve efficiency of the bargaining. This equilibrium also approximates the ex-post efficient equilibria as $\alpha > 1$ tends to 1. Consider the following strategy profile with $0 < \theta_1 < \theta_2 < 1$: the principal makes a price offer $\alpha \cdot b^2$ regardless of the message she has received and takes an action $y = \bar{y}^P(0, \theta_1) = \frac{\theta_1}{2}$ if the offer is rejected. Agent type in $[\theta_2, 1]$ fully reveals his private information by sending a truth-telling message $m = \theta$ in the cheap talk stage and accepts any offer less than or equal to $\alpha \cdot b^2$ with probability one but rejects any offer strictly greater than $\alpha \cdot b^2$ with probability one. Any agent type θ in $[0, \theta_2)$ uniformly randomizes over $M \setminus [\theta_2, 1]$ in the cheap talk stage and rejects the price offer $\alpha \cdot l(b)$ with probability one if $\theta \in [0, \theta_1]$ and accepts it with probability one if $\theta \in (\theta_1, \theta_2)$. If the principal makes any price offer $p \neq \alpha \cdot b^2$ and is rejected she takes action $y = m$ following the message from the agent in the cheap talk stage only if $m \in [\theta_2, 1]$. If the principal makes any price offer $p \neq \alpha \cdot b^2$ and is rejected but $m \notin [\theta_2, 1]$, then she takes action $y = \bar{y}^P(0, \theta_1) = \frac{\theta_1}{2}$.

Given the uniform prior, the decision rule of the agent of type $\theta \in [0, \theta_2]$ who partially babbles in the cheap talk stage should satisfy the *monotonicity*. Let $\theta_p \in [0, \theta_2]$ denote the agent type who is indifferent between accepting and rejecting the price offer p , if it exists. From the indifferent condition at θ_p , we have $\theta_p = 2(\sqrt{p} - b)$. Thus, the agent type $\theta \in [0, \theta_2]$ accepts price offer p with probability one if $\theta > \theta_p$, and otherwise rejects it with probability one. It is routine to show that the proposed strategies are mutual best responses given some belief derived by Bayes' rule.

The outcome of this equilibrium is described in Figure 9. Note that if $4(\sqrt{\alpha} - 1)b \geq 1$ or equivalently $\alpha \geq \left(1 + \frac{1}{4b}\right)^2$, this equilibrium boils down to the unique equilibrium of tacit bargaining. Similarly, when α goes to 1 this equilibrium converges to the ex-post efficient equilibrium of explicit bargaining. More importantly, as long as $\alpha < \left(1 + \frac{1}{4b}\right)^2$, communication can improve the efficiency of bargaining by expanding the set of agent types who would accept the equilibrium price offer from the principal. This result is summarized in the following proposition.

Proposition 6. *Assume the prior is uniform and utility functions are as in (11). For any given value $b > 0$, there exists an equilibrium in which pre-bargaining communication is used to improve the efficiency of bargaining by reducing the principal's inefficient screening of agent types if $\alpha < \left(1 + \frac{1}{4b}\right)^2$.*

Proof. See Appendix B. □

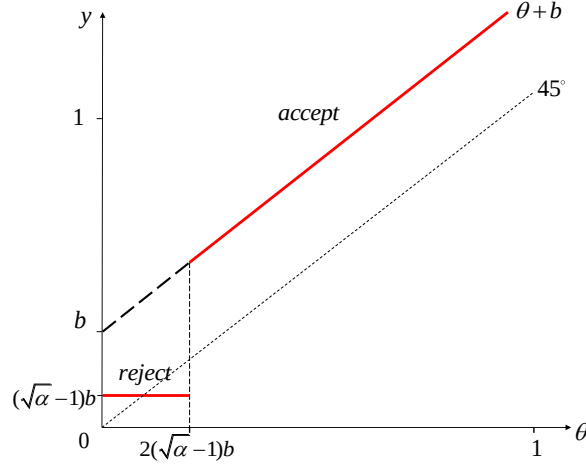


Figure 9: Equilibrium Outcome when $\alpha > 1$

5.4 Optimal Bargaining Mechanism

One important observation is that bargaining over decision rights considered in the previous sections is not the only bargaining mechanism for decision rights in our environment. For example, after the principal's price offer is rejected, the principal may want to make another price offer instead of taking action by herself. Alternatively, the informed agent may have bargaining power and make a take-it-or-leave-it price offer, as described by Lim [42]. It may be possible for the agent to make a price offer to buy decision rights after he rejects a price offer from the principal in the first round. We use the term "bargaining mechanism" to refer to any scheme through which principal and agent make offers directly, indirectly, once, repeatedly, sequentially, simultaneously, alternatively, and so on. Accordingly, we have *infinitely many* alternative bargaining mechanisms for decision rights. In this section, we follow a mechanism design approach (Myerson [47]) to demonstrate that if the players' utility function is quadratic or prior belief is uniform then the explicit bargaining we have considered is the optimal mechanism among the feasible mechanisms, as it achieves an upper bound of *ex-ante* social welfare.

A *bargaining mechanism* is one in which the informed agent sends a message to a mediator, who then credibly commits to a final allocation of decision rights and monetary transfers. By invoking the revelation principle, we can restrict our attention to *direct bargaining mechanisms* in which the informed agent reports a true state to a mediator without loss of generality. A direct mechanism is characterized by two outcome functions, denoted by $x(\cdot)$ and $p(\cdot)$, where $x(\theta)$ is the probability that the decision-right is transferred to the agent and $p(\theta)$ is the conditional payment from the agent to the principal when θ is the state reported by the

agent.¹⁶

Our goal is to find a mechanism that maximizes the social welfare, which is defined as the sum of the expected utilities of the principal and the agent.¹⁷ Given a mechanism (p, x) , define the *ex-ante* social welfare, which is denoted by $\mathcal{U}$, as follows.

$$\begin{aligned}
\mathcal{U} &= EU^A + EU^P \\
&= \int_{\Theta} \{x(\theta) \cdot U^A(\bar{y}^A(\theta), \theta, b) + (1 - x(\theta)) \cdot U^A(\bar{y}^P(p(\theta)), \theta, b) - p(\theta)\} f(\theta) d\theta \\
&\quad + \int_{\Theta} \{x(\theta) \cdot U^P(\bar{y}^A(\theta), \theta) + (1 - x(\theta)) \cdot U^P(\bar{y}^P(p(\theta)), \theta) + p(\theta)\} f(\theta) d\theta \\
&= \int_{\Theta} \{-x(\theta) \cdot l(b) + (1 - x(\theta)) \cdot [U^P(\bar{y}^P(p(\theta)), \theta) + U^A(\bar{y}^P(p(\theta)), \theta, b)]\} f(\theta) d\theta. \quad (12)
\end{aligned}$$

We state that a mechanism (p, x) is optimal if it maximizes the *ex-ante* social welfare. That is, an optimal mechanism is a solution to the following optimization problem:

$$\max_{p(\cdot), x(\cdot)} \mathcal{U}$$

subject to incentive compatibility constraints.

Equation (12) shows that an optimal mechanism should assign decision rights to the principal if and only if the sum of the interim utilities resulting from the action $\bar{y}^P(p(\theta))$ is greater than $-l(b)$, where $\bar{y}^P(p(\theta))$ is an action taken by the principal, who updates her belief after observing $p(\theta)$. Therefore, if a mechanism achieves a social welfare $\mathcal{U} > -l(b)$, there exists a nonempty set of agent types, denoted by S , such that

$$x(\theta) = \bar{x} \quad \text{and} \quad p(\theta) = \bar{p}, \quad \forall \theta \in S \quad (13)$$

and

$$\int_S [U^P(\bar{y}^P(\bar{p}), \theta) + U^A(\bar{y}^P(\bar{p}), \theta, b)] g(\theta) d\theta > -l(b) \quad (14)$$

¹⁶A more general class of mediation mechanisms would have the agent sending a report to a mediator, who then assigns decision rights to either the principal or agent according to a stochastic rule. If the agent receives decision rights, he pays an amount established by the mechanism, while if the principal retains authority, the mediator sends a recommended action. Note that in the latter case $p(\theta)$ itself can be interpreted and understood as a “recommendation” or a “message” from the mediator because upon rejection $p(\theta)$ becomes payoff-irrelevant. Therefore, we do not consider any additional recommendations by the mediator.

¹⁷In our environment, the principal’s optimal mechanism can easily be characterized. For example, assuming the participation constraint that the agent’s ex-ante or interim payoff must be greater than or equal to a particular value $\underline{U}^A$, a direct mechanism in which $x(\theta) = 1$ and $p(\theta) = \underline{U}^A$ for any $\theta \in [0, 1]$ that leads the agent to have the lowest possible payoff is incentive-compatible and individually rational.

where

$$\begin{aligned}\bar{y}^P(\bar{p}) &= \operatorname{argmax}_y \int_{\Theta} U^P(y, \theta) g(\theta) d\theta \\ &= \operatorname{argmax}_y \int_S U^P(y, \theta) f(\theta) d\theta\end{aligned}\tag{15}$$

and

$$g(\theta) \equiv \frac{q(\bar{p}|\theta) f(\theta)}{\int_{\Theta} q(\bar{p}|\theta') f(\theta') d\theta'}.$$

By approximating $U^A(y, \theta, b)$ with a linear function with a slope $l'(b)$ at $\theta = y$ and utilizing the strict concavity of the utility functions, we obtain

$$U^A(y, \theta, b) \leq U^A(y, y, b) + (y - \theta) \cdot \frac{\partial U^A(y, \theta, b)}{\partial \theta} \Big|_{\theta=y} = -l(b) + l'(b) \cdot (y - \theta) \tag{16}$$

for any $y \in Y$ and $\theta \in \Theta$. Taking integral in both sides, for any nonempty $S \subseteq \Theta$, we have

$$\int_S U^A(\bar{y}^P(\bar{p}), \theta, b) g(\theta) d\theta \leq -l(b) + l'(b) \int_S (\bar{y}^P(\bar{p}) - \theta) g(\theta) d\theta. \tag{17}$$

First, when the utility is quadratic, the second term of the right-hand side of Equation (17) becomes $-2b \int_S (\bar{y}^P(\bar{p}) - \theta) g(\theta) d\theta$, which equals zero according to the first-order condition on $\bar{y}^P(\bar{p})$ (the optimal action of the principal is an unbiased estimate of the state in equilibrium) so that the upper bound of the *ex-ante* social welfare is $-l(b)$. Second, when f is uniform, a similar analysis can be performed. From the first-order condition with a uniform prior, we obtain

$$\int_S U_1^P(\bar{y}^P(\bar{p}), \theta) d\theta = \int_{S_1} U_1^P(\bar{y}^P(\bar{p}), \theta) d\theta + \int_{S_2} U_1^P(\bar{y}^P(\bar{p}), \theta) d\theta = 0 \tag{18}$$

where $S_1 \equiv \{\theta \in S | \theta \geq \bar{y}^P(\bar{p})\}$ and $S_2 \equiv S \setminus S_1$. Since $U^P(y, \theta)$ is symmetric and $U_1^P(y, \theta)$ is strictly decreasing in θ , for any θ_1 and θ_2 , $|y - \theta_1| = |y - \theta_2|$ if and only if $U_1^P(y, \theta_1) = -U_1^P(y, \theta_2)$. Thus, the condition (18) is equivalent to $\int_S (\bar{y}^P(\bar{p}) - \theta) d\theta = 0$. Equation (17) then becomes

$$\int_S U^A(\bar{y}^P(\bar{p}), \theta, b) g(\theta) d\theta \leq -l(b) \tag{19}$$

which means that the upper bound of the *ex-ante* social welfare is $-l(b)$. This result is summarized in the following proposition.

Proposition 7. *Suppose that either utility is quadratic or prior f is uniform. Explicit bargaining is then optimal.*

The intuition behind this result is straightforward. A bargaining mechanism determines the final allocation of decision rights but has no effect on incentives in decision-making. That is, the final decision depends only on the decision-making party's own interest and the private

information that the party possesses. Crawford and Sobel [16] and Goltsman, *et al.* [24] show that more precise information is always beneficial *ex-ante* to both the principal and the agent when the utility is quadratic. For more general utility functions, the symmetry of the prior f for every subset of Θ is sufficient to obtain the same result. Therefore, the social welfare cannot be higher than $-l(b)$, which is the social welfare that results from the most informative decision-making.

Note that the efficiencies of the bargaining mechanisms are bounded away from the best efficiency with *ex-ante* social welfare $-\frac{b^2}{2}$ that could be achieved if the midpoint action between θ and $\theta + b$ is taken in every $\theta \in \Theta$ in the uniform-quadratic environment. This observation demonstrates that the noncontractibility of actions incurs a positive social cost.

6 Concluding Remarks

This paper studies bargaining over decision-making rights between an informed but self-interested agent and an uninformed principal. In the studied bargaining scheme the uninformed principal makes a price offer to the agent, who then decides to either accept or reject it. We show that if parties can communicate explicitly via cheap talk, a perfect Bayesian equilibrium exists in which pre-bargaining communication is used to achieve the ex-post efficiency by preventing the principal from inefficiently screening agent types, even when the conflict of interest is arbitrarily large. We also follow a mechanism design approach to demonstrate that under certain conditions, explicit bargaining is the optimal bargaining mechanism, i.e., the mechanism that maximizes the joint surplus of the parties.

There are several concluding remarks to be made. First, it is interesting to consider whether it would be possible to obtain the ex-post efficiency when the principal demands a price before the agent has observed his type. In this case, the agent's net value from rejecting a price offer p is $p - \int_{\Theta} U^A(y^*, \theta, b) f(\theta) d\theta$, where y^* is the *ex-ante* ideal action of the principal. Simultaneously, the net value of making an acceptable price offer p to the principal is $p - l(b) - \int_{\Theta} U^P(y^*, \theta) f(\theta) d\theta$. As a result, a price offer exists that leads to the efficient allocation of decision rights if and only if $\int_{\Theta} U^A(y^*, \theta, b) f(\theta) d\theta \geq \int_{\Theta} U^P(y^*, \theta) f(\theta) d\theta + l(b)$, or equivalently $\int_{\Theta} (U^A(y^*, \theta, b) + U^P(y^*, \theta)) f(\theta) d\theta \leq -l(b)$. These conditions will be satisfied if either f is uniform or U is quadratic.

Second, our main result, the existence of the ex-post efficient equilibrium, is sensitive to the principal having a small additional value to retain authority as discussed in Section 5.3. In particular, after the full transmission of information, the principal should want to take her ideal action herself rather than transfer authority when she would gain positive value from retaining authority. Indeed, Fehr, Herz and Wilkening [22] show that individuals exhibit

a strong tendency to retain authority even when its delegation is in their material interest, which suggests that they value authority *per se*. However, the existence of a ex-post efficient equilibrium in our model is restored if the agent also places the same value on obtaining authority. In this case, the agent is willing to accept a price offer that is slightly higher than the original equilibrium price offer $l(b)$ to compensate the principal for the loss of authority.

The allocation of decision rights through more complicated bargaining protocols such as alternative-offer bargaining would certainly extend the scope of our understanding of bargaining over decision rights. These issues represent promising subjects for future research.

References

- [1] Alonso, Ricardo, and Niko Matouschek. 2008. "Optimal Delegation." *The Review of Economic Studies*, 75: 259-293.
- [2] Aghion, Philippe, Mathias Dewatripont, and Patrick Rey. 2004. "Transferable control." *Journal of the European Economic Association*, 2: 115-138.
- [3] Aghion, Philippe, Mathias Dewatripont, and Patrick Rey. 2002. "On Partial Contracting." *European Economic Review*, 46: 745-753.
- [4] Aghion, Philippe, Jean Tirole. 1997. "Formal and Real Authority in Organizations." *The Journal of Political Economy*, 105: 1-29
- [5] Ambrus, Attila, and Georgy Egorov. 2009. "Delegation and Nonmonetary Incentives." *Harvard University*, Unpublished manuscript.
- [6] Ambrus, Attila, and Satoru Takahashi. 2008. "Multi-sender cheap talk with restricted state spaces." *Theoretical Economics*, 3(1): 1-27.
- [7] Arrunada, Benito, Luis Garicano, and Luis Vazquez. 2001. "Contractual Allocation of Decision Rights and Incentives: The Case of Automobile Distribution," *Journal of Law, Economics and Organization*, 17: 257-284.
- [8] Baker, George, Gibbons Robert, and Kevin J. Murphy. 2006. "Contracting for Control." Unpublished manuscript.
- [9] Battaglini, Marco. 2002. "Multiple Referrals and Multidimensional Cheap Talk." *Econometrica*, 70(4): 1379-1401.
- [10] Bester, Helmut. 2009. "Externalities, Communication and the Allocation of decision rights." *Economic Theory*, 41: 269-296.

- [11] Blume, Andreas, Oliver Board, and Kohei Kawamura. 2007. "Noisy talk," *Theoretical Economics*, 2: 395-440
- [12] Chatterjee, Kalyan, and William Samuelson. 1983. "Bargaining under incomplete information." *Operation Researches*, 31: 835-851.
- [13] Chen, Ying, Navin Kartik, and Joel Sobel. 2008. "Selecting Cheap-talk Equilibria" *Econometrica*, 76(1): 117-136.
- [14] Cho, In-Koo and David M. Kreps. 1987. "Signalling games and stable equilibria" *Quarterly Journal of Economics*, 102: 179-221.
- [15] Crawford, Vincent P. 1990. "Explicit Communication and Bargaining Outcome." *American Economic Review*, Papers and Proceedings of the Hundred and Second Annual Meeting of the American Economic Association, 80(2): 213-219.
- [16] Crawford, Vincent P. and Joel Sobel. 1982. "Strategic Information Transmission." *Econometrica*, 50: 1431-1451.
- [17] Dessein, Wouter. 2002. "Authority and Communication in Organizations." *Review of Economic Studies*, 69: 811-838
- [18] Farrell, Joseph. 1993. "Meaning and Credibility in Cheap Talk Games." *Games and Economics Behavior*, 5: 514-531.
- [19] Farrell, Joseph and Robert Gibbons. 1989. "Cheap Talk Can Matter in Bargaining." *Journal of Economic Theory*, 48: 221-237.
- [20] Farrell, Joseph and Robert Gibbons. 1989. "Cheap Talk with Two Audiences." *The American Economic Review*, 79(5): 1214-1223.
- [21] Farrell, Joseph and Matthew Rabin. 1996. "Cheap Talk." *The Journal of Economic Perspectives*, 10(3): 103-118.
- [22] Fehr, Ernst, Holger Herz, and Tom Wilkening. 2013. "The Lure of Authority: Motivation and Incentive Effects of Power." *American Economic Review*, 103(4): 1325-1359.
- [23] Gertner, Robert, Robert Gibbons, and David Scharfstein. 1988. "Simultaneous Signalling to the Capital and Product Markets." *The RAND Journal of Economics*, 19: 173-190.
- [24] Goltsman, Maria, Johannes Hörner, Gregory Pavlov and Francesco Squintani. 2009. "Arbitration, Mediation and Cheap Talk." *Journal of Economic Theory*, 144: 1397-1420.

- [25] Grossman, Sanford J., and Oliver Hart, 1986. "The Costs and Benefits of Ownership: A Theory of Vertical and Lateral Ownership." *Journal of Political Economy*, 94, 691-719.
- [26] Grossman, Stanford J., and Motty Perry. 1986. "Perfect Sequential Equilibrium." *Journal of Economic Theory*, 39: 97-119.
- [27] Grossman, Stanford J., and Motty Perry. 1986. "Sequential Bargaining under Asymmetric Information." *Journal of Economic Theory*, 39: 120-154.
- [28] Harrington, Joseph E. 1993. "The Impact of Reelection Pressure on the Fulfilment of Campaign Promises." *Games and Economic Behavior*, 5: 71-97.
- [29] Hart, Oliver, and Bengt Holmström. 2010. "A Theory of Firm Scope." *Quarterly Journal of Economics*, 125: 483-513.
- [30] Hart, Oliver, and John Moore. 1990. "Property Rights and the Nature of the Firm." *Journal of Political Economy*, 98: 1119-1158.
- [31] Holmström, Bengt. 1977. "On Incentives and Control in Organizations." Ph.D. Diss. Stanford University.
- [32] Holmström, Bengt. 1984. "On the Theory of Delegation." in M. Boyer and R. Kihlstrom (eds.) *Bayesian Models in Economic Theory*, New York: North-Holland.
- [33] Ivanov, Maxim. 2010. "Communication via a strategic mediator." *Journal of Economic Theory*, 145: 869-884.
- [34] Kohlberg, Elon and Jean-Francois Mertens. 1986. "On the Strategic Stability of Equilibria." *Econometrica*, 54: 1003-1037.
- [35] Kováč, Eugen and Tymofiy Mylovanov. 2009. "Stochastic Mechanisms in Settings without Monetary Transfers: Regular Case." *Journal of Economic Theory*, 144: 1373-1395.
- [36] Krähmer, Daniel. 2006. "Message-contingent Delegation." *Journal of Economic Behavior and Organization*, 60: 490-506.
- [37] Kreps, Daniel, and Robert Wilson. 1982. "Reputation and Imperfect Information." *Journal of Economic Theory*, 27: 253-180.
- [38] Kreps, Daniel, and Robert Wilson. 1982. "Sequential equilibrium." *Econometrica*, 50: 863-894.

- [39] Krishna, Vijay and John Morgan. 2004. "The art of conversation: eliciting information from experts through multi-stage communication," *Journal of Economic Theory*, 117: 147-179.
- [40] Krishna, Vijay and John Morgan. 2008. "Contracting for Information under Imperfect Commitment." *The RAND Journal of Economics*, 39: 905-925.
- [41] Lerner, Josh, and Robert P. Merges. 1998. "The Control of Technology Alliances: An Empirical Analysis of the Biotechnology Industry," *Journal of Industrial Economics*, 46: 125-156.
- [42] Lim, Wooyoung. 2012. "Selling Authority." *Journal of Economic Behavior & Organization*, 84: 393-415
- [43] Madrigal, Vincente, Tommy. C. C. Tan, and Sérgio Ribeiro da Costa Werlang. 1987. "Support Restrictions and Sequential Equilibria." *Journal of Economic Theory* 43: 329-334.
- [44] Matthews, Steven A. 1989. "Veto Threats: Rhetoric in a Bargaining Game." *Quarterly Journal of Economics*, 104: 347-369.
- [45] Matthews, Steven A., Masahiro Ocuno-Fujiwara and Andrew Postlewaite. 1991. "Refining Cheap-Talk Equilibria." *Journal of Economic Theory*, 55: 247-273.
- [46] Melumad, Nahum D. and Toshiyuki Shibano. 1991. "Communication in setting with no transfers." *The RAND Journal of Economics*, 22: 173-198.
- [47] Myerson, Roger. 1981. "Optimal Auction Design." *Mathematics of Operations Research*, 6: 58-73.
- [48] Myerson, Roger. 1978. "Refinements of the Nash equilibrium concept." *International Journal of Game Theory*, 15:133-154.
- [49] Nöldeke, Georg, and Eric van Damme. 1990. "Switching Away From Probability One Beliefs." SFB Discussion Paper No. A-304. *University of Bonn*, Unpublished Manuscript.
- [50] Rubinstein, Ariel. 1985. "A bargaining model with incomplete information about preferences." *Econometrica*, 53: 1151-1173.
- [51] Selten, Reinhard. 1975. "A reexamination of the perfectness concept for equilibrium points in extensive games." *International Journal of Game Theory*, 4: 25-55.

- [52] Sobel, Joel. 2013. “Giving and Receiving Advice.” *Advances in Economics and Econometrics* D. Acemoglu, M. Arellano, and E. Dekel (eds.).

Appendix A. Robustness of Ex-post Efficient Equilibria

In this appendix, we discuss the robustness of the ex-post efficient equilibrium. Given that the explicit bargaining game is not purely a cheap talk game, standard signaling refinements based on strategic stability (Kohlberg and Mertens [34]) may be effective. We first apply the Intuitive Criterion presented by Cho and Kreps [14].

In the ex-post efficient equilibrium, the equilibrium payoff for any agent type $\theta \in [0, 1]$ is $-l(b)$. Consider a deviation a_1 by agent type θ . First, if the considered deviation a_1 includes a decision rule that accepts price offers strictly higher than $l(b)$, the principal's best response is to make one of these acceptable price offers regardless of her belief. This best response by the principal results in a payoff strictly lower than $-l(b)$ for any agent type $\theta \in [0, 1]$. Second, if the considered deviation a_1 includes a decision rule that rejects any price offer strictly higher than $l(b)$, making an unacceptable price offer and taking action $\theta + b$ is the best response of the principal when she believes that agent type is $\theta + b$ with probability one. This best response results in a payoff of 0 (strictly bigger than $-l(b)$) for any agent type $\theta \in [0, 1]$. Thus, the “Equilibrium Domination Test” has no power. Furthermore, in the second deviation case, if we adopt the belief of the principal that assigns probability one to $\theta = 0$ (we are free to adopt any belief because the equilibrium domination test does not eliminate any θ), her best response is to make an unacceptable price offer and take action $y = 0$. This best response by the principal presents a payoff of $-l(\theta + b) \leq -l(b)$ for any $\theta \in [0, 1]$. Therefore, the ex-post efficient equilibrium satisfies the Intuitive Criterion. More generally, the efficient equilibrium survives any signaling refinement that imposes additional restrictions on the principal's belief off the equilibrium path because there are no unsent messages and both “accept” and “reject” are equilibrium-path events in the equilibrium. The only possible off-the-equilibrium path event is that a price $p \neq l(b)$ is set. However, this off-the-equilibrium state could be reached through a deviation by the principal instead of by the agent and thus plays no role in restricting the beliefs of the principal.

In the following sections, we turn our attention to two equilibrium refinements of cheap talk games using the concept of neologism-proofness developed by Farrell [18]¹⁸ and the NITS condition (no incentive to separate) developed by Chen, Kartik, and Sobel [13]. We will show that the efficient equilibrium satisfies both neologism-proofness and the NITS condition for any $b > 0$, whereas the babbling equilibrium does not satisfy either of them for certain values of b . We will further discuss the robustness of the ex-post efficient equilibrium against support restrictions, an assumption adopted in several studies such as those of Grossman and Perry

¹⁸Gertner, Gibbons, and Scharfstein [23], Farrell and Gibbons [19][20], and Matthews [44] are examples of papers that use neologism-proofness to refine their equilibrium outcomes.

[26][27], Harrington [28], Kreps and Wilson [37] and Rubinstein [50]. The extensive-form trembling-hand perfection used by Selten [51] and the sequential equilibria used by Kreps and Wilson [38] will also be discussed.

A.1. Neologism-proofness

Following Farrell [18], assume that for every nonempty subset X of Θ and for every perfect Bayesian equilibrium of the game, there exists a message $m(X)$ that is unused in the equilibrium and whose literal meaning is that $\theta \in X$. If the principal observes $m(X)$, she hypothesizes that some members of the specified subset X are responsible for the message and makes a price offer that is a best response given the posterior belief derived using Bayes' rule and her prior. For example, our analysis of tacit bargaining suggests that for any convex $X \subseteq \Theta$, the principal's best response upon receiving the message $m(X)$ is to make a price offer that is rejected by agent types in the lower half of X and accepted by the remaining agent types in X . The parties' behaviors in the remainder of the game satisfy sequential rationality, and the agent types' final payoffs from sending $m(X)$ are determined. Let $P(X)$ denote the set of all agent types who strictly prefer their payoffs from sending the message $m(X)$ to their equilibrium payoffs. We say that a subset X is *self-signaling* if $P(X) = X$. The neologism $m(X)$ is credible if X is self-signaling. If a credible neologism is available in an equilibrium, we say that such an equilibrium is not *neologism proof*.

The notion of neologism-proofness has a refining power in our game. To demonstrate this, suppose that the parties' prior beliefs are *uniform* over $[0, 1]$ and that their utilities are quadratic. In a babbling equilibrium, in which all agent types completely randomize over the message space and no information is transmitted, it might be the case that some agent types prefer to distinguish their types from others either because the unique action induced in equilibrium is too small ($\frac{1}{4}$) for them or the amount of money they must pay to the principal ($p = (\frac{1}{4} + b)^2$) is too high. Therefore, we investigate whether there is a self-signaling subset of the form $X = [\tilde{\theta}, 1]$ with $\tilde{\theta} > 0$. Suppose that $X = [\tilde{\theta}, 1]$ send a neologism to the principal. Then, according to Lemma 1 and the analysis of the tacit bargaining, the principal's optimal response is to make a price offer, p' that is rejected by $[\tilde{\theta}, \frac{\tilde{\theta}+1}{2}]$ but accepted by $(\frac{\tilde{\theta}+1}{2}, 1]$. From the indifference condition at $\frac{\tilde{\theta}+1}{2}$, we have $p' = (\frac{1-\tilde{\theta}}{4} + b)^2$. For X to be self-signaling, it is necessary and, for $\tilde{\theta} \in (0, 1)$, sufficient that i) the agent type $\tilde{\theta}$ is indifferent between sending the neologism that induces the action $\frac{3\tilde{\theta}+1}{4}$ and sending his equilibrium message that induces the action $\frac{1}{4}$ and ii) the agent type 0 does not want to deviate to the neologism $m(X)$. This requires that

$$\frac{1}{4} - \tilde{\theta} - b = -\frac{3\tilde{\theta}+1}{4} + \tilde{\theta} + b \quad (20)$$

and

$$-(\frac{1}{4} - b)^2 \geq -p' = -(\frac{1 - \tilde{\theta}}{4} + b)^2. \quad (21)$$

If equation (20) results in a value of $\tilde{\theta}$ in the range of $(0, 1)$ and the value of $\tilde{\theta}$ satisfies the inequality (21), then we have constructed a self-signaling subset X . It is immediately clear that $\tilde{\theta}$ satisfies both conditions if and only if $\frac{1}{24} \leq b < \frac{1}{4}$. Therefore, if $\frac{1}{24} \leq b < \frac{1}{4}$, any babbling equilibrium is not neologism proof.

It is well known that a neologism-proof equilibrium may select only a pooling equilibrium (Gertner, Gibbons, and Scharfstein [23]). Moreover, there may be no neologism-proof equilibrium (Matthews [44]). However, the ex-post efficient equilibrium is always neologism proof in our game.

Proposition 8. *For any $b > 0$, the ex-post efficient equilibrium is neologism proof.*

Proof. See Appendix B. □

The proof demonstrates that for *any* nonempty subset X of Θ there must be an agent type $\theta \in X$ such that sending the neologism $m(X)$ generates a payoff strictly less than $-l(b)$, the payoff from the ex-post efficient equilibrium. This result implies that the ex-post efficient equilibrium is robust to some variations of neologism-proofness, including the (*strongly / weakly*) *announcement-proofness* discussed by Matthews, Okuno-Fujiwara, and Postlewaite [45], who propose the additional requirement that *credible announcements* destroy the original equilibrium, where the *announcement* is a generalization of the neologism that allows deviant types to distinguish among themselves.

A.2. NITS (No Incentive To Separate)

Chen, Kartik, and Sobel [13] propose a criterion to select equilibria in Crawford and Sobel [16]: NITS, which stands for *no incentive to separate*. An equilibrium satisfies NITS if the agent of the lowest type weakly prefers the equilibrium outcome to credibly revealing his type. It has been demonstrated that equilibria satisfying NITS always exist in Crawford and Sobel [16], and the most informative equilibrium outcome is the unique equilibrium that satisfies NITS under the monotonicity condition M in Crawford and Sobel [16]. In this section, we apply NITS to our model and show that the criterion is selective; given some values of b , the babbling equilibrium does not survive, while the ex-post efficient equilibrium does survive for any $b > 0$.

Suppose that the parties' prior beliefs are *uniform* over $[0, 1]$ and that their utilities are quadratic. Note that in the babbling equilibrium, the agent of the lowest type receives $-(\frac{1}{4} -$

$0 - b)^2$. Thus, the babbling equilibrium does not satisfy NITS if and only if

$$-(\frac{1}{4} - b)^2 < -b^2,$$

which is equivalent to $b < \frac{1}{8}$.

It is obvious that NITS holds in the ex-post efficient equilibrium in which all agent types reveal their types truthfully.

Proposition 9. *For any $b > 0$, the ex-post efficient equilibrium satisfies NITS.*

A.3. Support Restriction, Sequential Equilibria and Trembling-hand Perfection

One might be tempted to argue that the ex-post efficient equilibrium is not particularly reasonable because it may not satisfy the “support restriction”, the assumption adopted by several papers such as those of Grossman and Perry [26][27], Harrington [28], Kreps and Wilson [37] and Rubinstein [50]. This restriction requires that the support of the beliefs at an information set should be contained in the supports of the beliefs at preceding information sets.

In our ex-post efficient equilibrium, the principal has probability-one beliefs after receiving messages in the cheap talk stage and switches away from these beliefs to a new belief that assigns probability one to the type $\theta = 0$ after observing the “rejection” of the equilibrium price offer; therefore, the equilibrium may violate the support restriction. There are two possible responses. First, the “rejection” of $p = l(b)$ is not an off-the-equilibrium-path event because agent type $\theta = 0$ rejects it with positive probability. That is, the belief assignment after the rejection is not a consequence of changing beliefs but the consequence of a Bayesian update. Second, it has been shown that the support restriction may be based on erroneous interpretation of the concept of a belief in some games and that furthermore, violations of the support restriction may represent a sensible reasoning process that supports interesting equilibrium behaviors, as described by Madrigal, Tan and Werlang [43] and Nöldeke and van Damme [49]:

... violations of the support restriction may very well reflect the fact that once a deviation from equilibrium behavior has been observed, a reassessment of all previous beliefs - which were based on the assumption that equilibrium strategies are followed - is called for. In this light such “switching beliefs” is not an unfortunate problem that , which cannot be avoided in some cases, but actually is a natural consequence of observing a deviation. (Nöldeke and van Damme [49], p. 9.)

One may also wonder if the ex-post efficient equilibrium is robust against the sequential rationality described by Kreps and Wilson [38] and the extensive form trembling-hand perfection discussed by Selten [51]. While the definition of sequential equilibria does not directly

apply to our game, the ex-post efficient equilibrium does not violate a sensible implementation of sequential equilibria. To see this, suppose that independent trembles occur at every decision node in the extensive form of our game. Due to these trembles, the principal makes some offers $p \neq l(b)$ with a positive probability. The resulting beliefs after every history can be derived using Bayes' rule, which implies that after a price offer is rejected, the principal maintains the belief she updated through the cheap talk stage. This demonstrates that the principal's beliefs at every decision node in the efficient equilibrium are consistent in the sense described by Kreps and Wilson [38].

We now shift our attention to the perfection of the equilibrium, which requires a mutual best response in addition to consistent beliefs in every perturbed game. It might be the case that a price offer larger than $l(b)$ would be a profitable deviation for the principal experiencing the trembles. There are two possible effects of trembles on the principal's incentives when choosing a price offer. First, compared to the original equilibrium path, the agent pays more when a price offer (larger than $l(b)$) is accepted. However, after the price offer is rejected, the action choices made by the principal are not as beneficial as selling the decision rights at price $l(b)$ because the information provided by the agent in the cheap talk stage includes noise from the trembles. The second effect must dominate the first if we consider a reasonable, fully mixed behavioral agent strategy in which a higher price offer is accepted with a smaller probability.¹⁹ A more serious problem emerges when we look at the agent's incentives: due to the trembles, the agent's message m induces the principal's action $y = m$ with positive probability. Taking this into account, agent type θ strictly prefers to send a message $\theta + b$ over a message θ . This demonstrates that our efficient equilibrium fails to satisfy the extensive form trembling-hand perfection. This failure stems primarily from the discontinuity in the agent's decision rule at $p = l(b)$.

Appendix B. Omitted Proofs

Proof of Lemma 1. First, any price offer $p < 0$ is accepted by all agent types because for any $\theta \in [0, 1]$, $U^A(y, \theta, b) < -p$ for any action $y \in \mathbb{R}$. Therefore, in the remainder of the proof, take $p \geq 0$. Let $\bar{y}$ denote the principal's action induced by the offer p . Suppose that an agent type $\bar{\theta} \in [0, 1]$ accepts the price offer p with positive probability in equilibrium. We then have $U^A(\bar{y}, \bar{\theta}, b) \leq -p$. By continuity, there exists a $\theta_p \in [0, 1]$ such that $U^A(\bar{y}, \theta_p, b) = -p$. (Otherwise, we have $U^A(\bar{y}, \theta, b) < -p$ for any $\theta \in [0, 1]$, which would imply that all agent types accept the offer and the proof is complete.) Now, suppose that $\theta_p > \bar{\theta}$. Then, by the concavity

¹⁹This means that expensive mistakes could be made less frequently, as considered in the discussion of properness in Myerson [48].

of U^A , the set of agent types who reject p with probability one is $(\theta_p, \theta']$ with $\theta_p < \theta' \leq 1$. Because the agent type θ' rejects the offer, we have $U^A(\bar{y}, \theta', b) \geq -p$. However, using **(B4)** and Bayes' rule, we have $\bar{y} = y(\theta_p, \theta')$, and by Condition 1, $U^A(\bar{y}, \theta', b) < U^A(\bar{y}, \theta_p, b) = -p$, which leads to a contradiction. Therefore, we have $\theta_p \leq \bar{\theta}$. Then, utilizing the strict concavity of U^A , we obtain $U^A(\bar{y}, \theta, b) < -p$ for all $\theta > \theta_p$. This implies that the monotonicity holds under Condition 1.

Proof of Proposition 2. Monotonicity implies that for any $p \in \mathbb{R}$, both $\Theta(p)$ and $\Theta'(p)$ are convex if they are nonempty. Further, $\Theta(p)$ cannot be to the left of $\Theta'(p)$. These guarantee that for any $p \in \mathbb{R}$, there is at most one agent type who is indifferent between accepting and rejecting the offer. Let $\theta_p \in [0, 1]$ denote this agent type if it exists. We can then write that $\Theta(p) = (\theta_p, 1]$ and $\Theta'(p) = [0, \theta_p)$. From the indifference condition at θ_p , we have

$$p = l(|y^P(p) - \theta_p - b|), \quad (22)$$

where

$$y^P(p) = \arg \max_y \int_0^{\theta_p} -l(|y - \theta|) \cdot \frac{f(\theta)}{F(\theta_p)} d\theta = y(0, \theta_p). \quad (23)$$

From (22), we have

$$\theta_p = y(0, \theta_p) + l^{-1}(p) - b. \quad (24)$$

The principal chooses p^* to solve

$$\begin{aligned} \max_{p \in \mathbb{R}} EU^P &= \int_0^{\theta_p} -l(|y^P(p) - \theta|) d\theta + (1 - \theta_p)(p - l(b)) \\ &\text{s.t. (24).} \end{aligned} \quad (25)$$

From (24), we have

$$\frac{\partial \theta_p}{\partial p} = \frac{\partial y^P(p)}{\partial \theta_p} \cdot \frac{\partial \theta_p}{\partial p} + \frac{\partial l^{-1}(p)}{\partial p}.$$

After some rearrangement, we obtain

$$\frac{\partial \theta_p}{\partial p} = \frac{1}{1 - \frac{\partial y^P(p)}{\partial \theta_p}} \cdot \frac{\partial l^{-1}(p)}{\partial p}.$$

Because $\frac{\partial l^{-1}(p)}{\partial p} \geq 0$, we have $\frac{\partial \theta_p}{\partial p} \geq 0$. This equation together with $0 < \frac{\partial y^P(p)}{\partial \theta_p} < 1$ implies that $\frac{\partial y^P(p)}{\partial p} > 0$. Taking the derivative of (25) w.r.t. p yields

$$\frac{\partial EU^P}{\partial p} = \frac{\partial \theta_p}{\partial p} \cdot (-l(|y(0, \theta_p) - \theta_p|) - p + l(b)) + (1 - \theta_p). \quad (26)$$

At $\theta_p = 0$, we have

$$\frac{\partial EU^P}{\partial p} \Big|_{\theta_p=0} = \frac{\partial \theta_p}{\partial p} \Big|_{\theta_p=0} \cdot (l(b) - l(b)) + (1 - 0) = 1 > 0. \quad (27)$$

At $\theta_p = 1$, we have

$$\frac{\partial EU^P}{\partial p} \Big|_{\theta_p=1} = \frac{\partial \theta_p}{\partial p} \Big|_{\theta_p=1} \cdot (-l(|y(0, 1) - 1|) - l(|1 - y(0, 1) + b|) + l(b)) < 0. \quad (28)$$

Therefore, by continuity, $\theta_{p^*} \in (0, 1)$. This completes our proof.

Proof of Proposition 3. The proof is constructive. Consider the following strategies and beliefs:

- i) The agent type θ fully reveals his private information by sending a message $m = \theta$.
- ii) For any $m \in M$, the principal makes the price offer $l(b)$.
- iii) For any $\theta \in [0, 1]$, the agent accepts the offer p with probability one if $p < l(b)$ but rejects p with probability one if $p > l(b)$, regardless of the message he sent. The agent type $\theta = 0$ randomizes between accepting and rejecting the offer $p = l(b)$ and any other agent types in $(0, 1]$ accept $l(b)$ with probability one.
- iv) If a price offer $p = l(b)$ is rejected, the principal takes an action $y = 0$ regardless of the message she received. If a price offer $p \neq l(b)$ is rejected, the principal who received the message m takes the action $y = m$.
- v) For any $m \in M$, $\rho_1(\theta|m) = \begin{cases} 0 & \forall \theta \in [0, 1] \setminus m, \\ 1 & \text{if } \theta = m. \end{cases}$
- vi) For any $m \in M$ and $p = l(b)$, $\rho_2(\theta|m, p) = \begin{cases} 0 & \forall \theta \in (0, 1], \\ 1 & \text{if } \theta = 0. \end{cases}$
- vii) For any $m \in M$ and any $p \neq l(b)$, $\rho_2(\theta|m, p) = \begin{cases} 0 & \forall \theta \in [0, 1] \setminus m, \\ 1 & \text{if } \theta = m. \end{cases}$

First, consider the agent's incentive. Given the principal's strategy and beliefs described above, the agent has no incentive to deviate in his message rule because the principal makes the message-independent price offer $l(b)$. For any $m \in M$, any agent type $\theta \in (0, 1]$ accepts an offer p with probability one if $p \leq l(b)$ because he receives $-p$, which is greater than or equal to $-l(b)$, from accepting the offer, but the expected payoff of rejecting the offer is

$$-l(|0 - \theta - b|) = -l(\theta + b) \leq -l(b), \quad \forall \theta \in [0, 1].$$

The previous calculation also shows that agent type $\theta = 0$ is indifferent between accepting $p = l(b)$ and rejecting it. Any agent type $\theta \in [0, 1]$ who fully reveals his private information rejects an offer p with probability one if $p > l(b)$ because he receives $-p$, which is less than $-l(b)$, for accepting the offer, but the expected payoff of rejecting the offer is $-l(b)$.

Second, consider the principal's incentive. Given the agent's strategy and beliefs above, the principal's optimal behavior after observing a truthful message (or a message θ) is to make the price offer $l(b)$, as any offer less than $l(b)$ will be accepted by all agent types with probability one and provide her with an expected payoff of strictly less than 0, the principal's expected payoff from making the offer $l(b)$. Any offer greater than $l(b)$ will be rejected with probability one and induce the action θ from the principal, which provides the principal with the expected payoff 0.

According to this construction, the price offer $l(b)$ is rejected with positive probability on the equilibrium path only if $\theta = 0$. Under Bayes' rule, the principal should assign probability one to $\theta = 0$ after the price offer is rejected. The principal's action rule described above is sequentially rational under the beliefs we have adopted. This completes our proof.

Proof of Proposition 4. Consider the strategy profile proposed above with an equilibrium path price offer $l(|b|)$ and a threat action z . It suffices to show that there exists a $z \in \Theta$ such that $-l(|b|) > -l(|z - \theta - b|)$ for all $\theta \in \Theta \setminus \{z\}$. Because $l'(\cdot) \geq 0$, the condition is equivalent to $|b| < |z - \theta - b|$ or

$$b \cdot b < b \cdot b + 2b \cdot (\theta - z) + (\theta - z) \cdot (\theta - z), \quad \forall \theta \in \Theta \setminus \{z\}.$$

Consider the function $b \cdot \theta$. Note that it is a continuous function on the compact set Θ and therefore achieves a minimum on Θ . Let z be the minimizer of $b \cdot \theta$ on Θ . Then 1) $b \cdot \theta \geq b \cdot z$ for all $\theta \in \Theta$ and 2) $(\theta - z) \cdot (\theta - z) > 0$ for all $\theta \in \Theta \setminus \{z\}$. This completes the proof.

Proof of Proposition 5. The case of $z = 1$ is omitted to avoid redundancy from our previous discussion of the case of $z = 0$. It remains to be shown that there is no continuous $b(\cdot)$ that satisfies the condition (10). Suppose that this is not the case. Take an arbitrary θ such that $\theta < \bar{\theta}$. Then $b(\theta) < 0$, because otherwise, $|b(\theta) + \theta - \bar{\theta}| < |b(\theta)|$, which contradicts (10) and the condition that $\bar{\theta}$ is a global maximizer of $|b(\cdot)|$. Similarly, take an arbitrary θ such that $\theta > \bar{\theta}$. Then $b(\theta) > 0$, because otherwise $|b(\theta) + \theta - \bar{\theta}| < |b(\theta)|$. The continuity of $b(\cdot)$ implies that $b(\bar{\theta}) = 0$, which contradicts the condition that $\bar{\theta}$ is a global maximizer of $|b(\cdot)|$.

Proof of Proposition 6. In order for the considered strategy profile to constitute equilibrium, we should have $\theta_2 = 2\theta_1$ and $\theta_1 = \theta_{p=\alpha \cdot b^2}$, or equivalently $\theta_1 = 2(\sqrt{\alpha} - 1)b$ and $\theta_2 = 4(\sqrt{\alpha} - 1)b$. A few things remain to be confirmed. First, no agent type has incentives to deviate in the cheap talk stage because the principal's price offer is message-independent. Second, the indifference condition at $\theta_1 = \theta_{p=\alpha \cdot b^2}$ together with the use of the strictly concave utility ensures that no agent type $\theta \in [0, 1]$ has an incentive to deviate in the bargaining

stage. The principal does not have an incentive to make any offer $p < \alpha \cdot b^2$ because 1) when $\theta \in [\theta_2, 1]$ such price offers will be accepted but provide a strictly lower payoff to the principal than making the equilibrium price offer and 2) when $\theta \in [0, \theta_2]$ making the price offer $\alpha \cdot b^2$ is already proven to be optimal in Section 3. Similarly, the principal does not have a strict incentive to make any offer $p > \alpha \cdot b^2$ because 1) when $\theta \in [\theta_2, 1]$ such price offers will be rejected, and the actions that would be taken by the principal after $p > \alpha \cdot b^2$ is rejected will provide the same payoff as could be obtained by making the offer $p = \alpha \cdot b^2$, and 2) when $\theta \in [0, \theta_2]$ making the price offer $\alpha \cdot b^2$ is already proven to be optimal. After the price offer $\alpha \cdot b^2$ is rejected, the principal believes that the true state of the world is in $[0, \theta_1]$ which is derived by Bayes' rule. And given this belief, taking action $\frac{\theta_1}{2} = (\sqrt{\alpha} - 1)b$ is sequentially rational for the principal. This completes our proof.

Proof of Proposition 8. Note that for any $\theta \in \Theta$, the agent's payoff in the truth-telling equilibrium is $-l(b)$. Any singleton subset of Θ could not be self-signaling because sending a neologism message by himself reveals his true type to the principal, so the agent type in the set will not receive more than $-l(b)$. Therefore, suppose that an arbitrary nonsingleton subset $\hat{\Theta}$ of Θ sends a neologism message $m(\hat{\Theta})$ to the principal. Let $\hat{p}$ denote the price offer induced by $\hat{\Theta}$. Then, $\hat{p} \geq l(b)$ because otherwise, all of the agent types in Θ would be strictly better off if they accepted the offer $\hat{p}$, which implies that $\hat{\Theta} = \Theta$. This result is contradictory, as Lemma 1 implies that making a price offer $\hat{p} < l(b)$ is never optimal for a principal who believes that $\hat{\Theta} = \Theta$. For the neologism $m(\hat{\Theta})$ to be credible, all of the agent types in $\hat{\Theta}$ should reject $\hat{p}$ because accepting $\hat{p}$ provides them with a maximum of $-l(b)$. However, rejecting $\hat{p}$ cannot provide all of the agent types in $\hat{\Theta}$ with more than $-l(b)$ because the principal's action induced by $\hat{\Theta}$, denoted by $y^P(\hat{p})$, always lies in the interior of $C(\hat{\Theta})$, which represents the convex hull of $\hat{\Theta}$. As a result, there is always a positive measure of agent types in the neighborhood of $y^P(\hat{p})$ who receive a strictly lower payoff than $-l(b)$. Therefore, $m(\hat{\Theta})$ cannot be a credible neologism for any $\hat{\Theta} \subseteq \Theta$.