

Complete and incomplete information bargaining in a model with independent private values*

Simon Loertscher[†]

Leslie M. Marx[‡]

January 30, 2026

Abstract

Bargaining is central to economics, yet subject to rich, often contradictory assumptions and results, apparently along a divide between complete and incomplete information settings. We analyze an independent private values model that, as a function of primitives, exhibits essential features of both complete and incomplete information bargaining. Whether the ownership structure and bargaining weights matter for ex post efficient bilateral bargaining depends on the extent to which type distributions overlap. Equal bargaining weights are always consistent with ex post efficiency whereas there is no ownership structure that always is. With overlapping supports, ex post efficiency requires equal bargaining weights.

Keywords: ex post efficiency, property rights, bargaining power, countervailing power, decreasing marginal values

JEL Classification: D44, D82, L41

*We are grateful for suggestions and feedback from Alessandro Bonatti, Piotr Dworczak, Will Fuchs, Yingni Guo, Martin Peitz, Patrick Rey, Vasiliki Skreta, Mia Tam, Steve Williams, and seminar participants at Australia National University, Duke University, MIT Sloan, Universitat Autònoma de Barcelona, Universitat Pompeu Fabra, University of Chicago, University of Melbourne, University of Southern California, University of Texas, Northwestern University, NYU Abu Dhabi, 2025 BEAT Trobada, 2023 CRESSE-Lingnan Competition Policy Workshop in Hong Kong, 2025 Australasian Economic Theory Workshop, 9th BEET Workshop, EARIE 2025, APIOC 2025, and SEA 2025. Financial support from the Australian Research Council Discovery Project Grants DP200103574 and DP250100407 is gratefully acknowledged.

[†]Department of Economics, Level 4, FBE Building, 111 Barry Street, University of Melbourne, Victoria 3010, Australia. Email: simonl@unimelb.edu.au.

[‡]Duke University, 100 Fuqua Drive, Durham, NC 27708, USA: Email: marx@duke.edu.

1 Introduction

Bargaining is central to economics, yet subject to rich, often contradictory assumptions and results. For example, the Coase Theorem states conditions under which the initial ownership structure is irrelevant for whether the final allocation is efficient. This Coasian irrelevance contrasts sharply with the available evidence of long-lasting effects of initial ownership.¹ Similarly, notwithstanding its popular appeal, John Kenneth Galbraith’s concept of countervailing power has been controversial since the beginning. While Galbraith found it to be of “substantial, and perhaps central, importance,” George Stigler lamented the lack of explanation for “why bilateral oligopoly should in general eliminate, and not merely redistribute, monopoly gains.”² The impossibility results of Vickrey and Myerson-Satterthwaite regarding ex post efficient trade with incomplete information seem difficult to reconcile with the efficiency imposed axiomatically by Nash and Shapley for complete information bargaining. This has led some authors to conclude that the quest for efficient bargaining is “fruitless.”³ The axiomatically imposed efficiency for complete information bargaining is also challenged by empirical evidence that bargaining breakdown is a salient feature of reality.⁴ In light of this, readers who wonder what lessons can be drawn from this seventy-five years old economics literature may perhaps be forgiven.

These divides seem largely driven by whether bargaining is modeled as exhibiting complete or incomplete information. For example, the literature with incomplete information in the form of independent private values has shown that both the ownership structure and bargaining weights matter for the size of the pie in, say, bilateral bargaining.⁵ However, this impression is not correct as we show in this paper. Specifically, we study an independent private values model and show that it exhibits essential features of both complete and incomplete information bargaining models. Even though information is always incomplete, depending on the parameterization, incomplete information bargaining can effectively look like (generalized) Nash bargaining. To the best of our knowledge, this is the first paper to provide an analysis of bargaining when bargaining weights and the ownership structure can vary simultaneously and affect, in general, whether bargaining is efficient and that encompasses central features of complete and incomplete information bargaining in a systematic and, with

¹Bleakley and Ferrie (2014) document that the effects of initial land ownership on the Georgia frontier only vanish after 150 years.

²See Galbraith (1954, p. 1) and Stigler (1954, p. 13). Steptoe (1993) summarizes its popular appeal by noting that the notion of buyer power—often taken to be synonymous with countervailing power—is sometimes embraced by courts as if it had “talismanic power.”

³Ausubel et al. (2002, p. 1934).

⁴Bargaining breakdown is observed 38% of the time in business-to-business used car negotiations and 45% of the time on eBay; see Larsen et al. (2024) and Backus et al. (2020).

⁵See e.g. Cramton et al. (1987) and Williams (1987).

the benefit of hindsight, intuitive way. In particular, it captures all of the above—Coasian irrelevance and non-Coasian relevance of the ownership structure, the central importance of Galbraith’s countervailing power and Stigler’s redistribution of monopoly gains, and efficient bargaining and inefficient bargaining breakdown—as special cases.

Varying the degree of overlap between the agents’ type distributions by allowing the support of some agents’ distributions to shift and taking the mechanism design approach to bargaining of Myerson and Satterthwaite (1983) that has proven productive for applied work,⁶ we model bargaining as a mechanism that maximizes the weighted sum of the agents’ expected payoffs, subject to incentive compatibility, interim individual rationality, and a no-deficit constraint for the mechanism designer. We show that, for bargaining to be ex post efficient, the ownership structure is less important than bargaining weights insofar as the ownership structure is relevant for a smaller subset of environments than are the bargaining weights. At the same time, getting the ownership structure right is, theoretically, more difficult because there is no single ownership structure that is always consistent with ex post efficient bargaining, whereas equal bargaining weights always are.⁷ To be more specific, with two agents and constant marginal values, ex post efficiency requires equal bargaining weights and appropriate ownership structure, which excludes extremal ownership, if the supports overlap. With nonoverlapping supports, the ownership structure is irrelevant for ex post efficiency because the ownership share of the weaker agent only affects the amount of the asset to be bargained over. In contrast, bargaining weights continue to matter for ex post efficiency with nonoverlapping supports as long as the exertion of monopoly or monopsony power is inefficient. However, the set of bargaining weights that induce ex post efficient bargaining increases as the stronger agent’s type distribution shifts to the right and eventually encompasses all bargaining weights. Consequently, from that point onward, neither bargaining weights nor ownership matters for ex post efficiency, so that bargaining resembles (generalized) Nash bargaining in which bargaining weights only affect the distribution of the gains from trade, even though information is incomplete.

The complete information approach to bargaining is sometimes deemed advantageous in applications because of (perceived) greater tractability. Researchers with that in mind may be inclined to defend this approach on grounds that, for the application at hand, there is little private information. Our framework provides a way to formalize and conceptualize the otherwise vague or vacuous notion of little private information.⁸ Private information is

⁶See, for example, Backus et al. (2019), Backus et al. (2020), Larsen (2021), Larsen et al. (2024), Larsen and Zhang (2025), Barkley et al. (2025a,b), and Backus et al. (forth.).

⁷Whether bargaining is ex post efficient depends simultaneously on the ownership shares and the bargaining weights. The more formal statement is that the intersection across all environments of the two-dimensional set of ownership shares and bargaining weights is the empty set paired with equal ownership.

⁸For example, the impossibility theorem of Myerson and Satterthwaite (1983) holds for any positive

negligible if the gap between the supports is so large that bargaining is efficient independent of bargaining weights and the ownership structure.⁹

Away from ex post efficiency, we characterize the Pareto frontiers of the agents' ex ante expected payoffs. For constant marginal values, we also show that the analysis extends directly to settings with many agents and, for the bilateral trade setting, we provide an indirect implementation via fee-setting mechanisms.

We show that the main insights with two agents generalize to settings with decreasing marginal values in which the agents' private information pertains to the intercepts of their linear inverse demand functions. There are, however, also subtle and surprising differences relative to constant marginal values. For example, with identical supports, there is only one ownership structure that is consistent with ex post efficient bargaining, whereas with constant marginal values there is a nondegenerate interval of ownership shares that are consistent with that. Similarly, with overlapping, nonidentical supports, the set of ownership structures consistent with ex post efficiency need not be convex, and extremal ownership can render bargaining ex post efficient even when supports overlap, contrasting with the impossibility result of Myerson-Satterthwaite for constant marginal values. Nonetheless, just as with constant marginal values, equal bargaining weights are necessary for ex post efficient bargaining as long as the supports overlap.

From the Coase Theorem (Coase, 1960) to the concept of countervailing power (Galbraith, 1952, 1954; Stigler, 1954) to the theory of the firm (Grossman and Hart, 1986; Hart and Moore, 1990) to merger cases (see, e.g., Lee et al., 2021) to empirical studies (Backus et al., 2020; Larsen, 2021; Byrne et al., 2022; Larsen et al., 2024; Larsen and Zhang, 2025; Barkley et al., 2025a,b; Backus et al., forth.), bargaining is central to economics.¹⁰ Our paper contributes to the literature on bargaining, including the strands of literature on com-

densities whose supports overlap, no matter how skewed these are.

⁹In practice, this would correspond to a situation in which, say, a downstream firm's value for a supplier's input is always well above the supplier's cost, and the downstream firm has no scope for internal production, which may be descriptive of the healthcare industry and insurers and hospitals. In settings in which bargaining breakdown is a real-world phenomenon, a model with incomplete information would seem more appropriate. See Ho and Lee (2017) for healthcare, Van der Maelen et al. (2017) for retail markets, Frieden et al. (2020) for media markets. A recent FCC proceeding focuses on cable and satellite TV blackouts.

¹⁰Avignon et al. (2025) and Demirer and Rubens (2025) examine an upstream monopsony and downstream monopoly engaged in generalized Nash bargaining over a linear wholesale price and show that a unique bargaining weight maximizes social surplus, which is in the same spirit as our countervailing power result for the case of overlapping supports. (Because of the exertion of market power on both the input and the output market, the unique social-surplus-maximizing bargaining weight of Avignon et al. (2025) and Demirer and Rubens (2025) may not be $1/2$ as in our setup with overlapping supports.) This similarity of findings for fairly different settings—incomplete information bargaining with independent private values for a fixed resource here versus generalized Nash bargaining in an oligopoly setting there—is reassuring in that it suggests that the results do not hinge on the specifics of the setups and may indeed reflect a deeper underlying force.

plete information bargaining in the tradition of Nash (1950) and Shapley (1951), and those on incomplete information bargaining along the lines of Vickrey (1961) and Myerson and Satterthwaite (1983), by providing and analyzing a framework that simultaneously and separately accounts for the effects of bargaining power and ownership on bargaining outcomes.¹¹ In particular, taking the same as-if approach as Myerson and Satterthwaite (1983), Ausubel et al. (2002), Loertscher and Marx (2022), and Choné et al. (2024) that models incomplete information bargaining as intermediated by a mechanism designer, we analyze a general partnership model (Cramton et al., 1987) that permits heterogeneous distributions and supports and unequal bargaining weights. In our framework, bargaining breakdown implies that bargaining is not ex post efficient. That is, observing bargaining leading to no trade means that the problem at hand is not such that trade always maximizes social surplus ex post.¹² This suggests that methods such as those developed by Barkley et al. (2025a) to empirically account for rejected offers may be particularly valuable for applications.

The paper contributes to the mechanism design literature that assumes constant marginal values by applying the methods that Myerson (1981), Myerson and Satterthwaite (1983), and Williams (1987) developed for settings with one-sided and two-sided private information to partnership models, which exhibit countervailing incentives, such as Lu and Robert (2001), Loertscher and Wasser (2019), Loertscher and Marx (2023), and Loertscher and Muir (2025).^{13,14} The partnership framework analyzed here is a generalization of the settings with one-sided and two-sided private information in Loertscher and Marx (2019, 2022) and the setting of Williams (1987).¹⁵

¹¹There is also a strand of literature on one-to-many bargaining—such as between a developer and land owners—with complete information in which the principal lacks commitment power; see, for example, Xiao (2018) or Uyanik and Yengin (2023).

¹²To see this, notice that with nonextremal ownership, the possibility that the efficient quantity traded is zero is a probability zero event. With constant marginal values, extremal ownership and overlapping supports, it is ex post efficient that the seller retains its full ownership if its value exceeds the buyers, but from Myerson and Satterthwaite (1983), we know that ex post efficient trade is impossible in this case.

¹³Focusing on ex post efficiency, Loertscher and Marx (2024) study variants of partnership models, sometimes also referred to as “asset market models,” in which agents’ types can be multi-dimensional, while Liu et al. (2026) provide an empirical application of an asset market model.

¹⁴Studying informed-principal problems, Mylovanov and Tröger (2014) solve for the mechanism that is optimal for one agent in a bilateral partnership problem. The mechanism that maximizes one agent’s expected payoff is encompassed as a special case of our analysis, which sidesteps the informed-principal aspect of the problem by assuming that the mechanism is designed and run by an intermediary. However, away from extremal bargaining weights, the informed-principal problem does not appear well defined.

¹⁵Our paper generalizes the prior literature with constant marginal values by not imposing equal bargaining weights, extremal ownership, identical distributions, and/or a common supports for type distributions. For example, Myerson and Satterthwaite (1983), Cramton et al. (1987), Gresik and Satterthwaite (1989), Makowski and Mezzetti (1993), Lu and Robert (2001), Che (2006), Figueroa and Skreta (2012), Loertscher and Wasser (2019), and Liu et al. (2026) assume equal weights. Williams (1987) and Loertscher and Marx (2022) allow for unequal bargaining weights but assume extremal resource ownership so that each agent is either ex ante known to be a buyer or a seller. Loertscher and Marx (2024) allow for nonextremal ownership

The paper also expands the mechanism design approach to incomplete information bargaining to settings with decreasing marginal values where the agents' private information pertains to the intercept of their linear inverse demand functions. Methodologically, we show that the underlying mechanics and many of the key results carry over from the framework with constant marginal values, but that there are also subtle and important differences.

The remainder of this paper is organized as follows. Section 2 lays out the setup. In Section 3, we derive the bargaining mechanism for constant and decreasing marginal values. Focusing on constant marginal values, in Section 4, we analyze ex post efficient bilateral bargaining, and in Section 5, we characterize the Pareto frontiers for bilateral bargaining. In Section 6, we discuss implementation, provide results for the case of more than two agents, and analyze ex post efficiency and Pareto frontiers for the case of decreasing marginal values. Section 7 concludes the paper.

2 Setup

In this section, we introduce the model, the bargaining mechanism, and objects of particular interest. The section concludes with a brief discussion of our modeling assumptions.

2.1 Environment

Let $\mathcal{N}$ with cardinality $n \geq 2$ denote the set of agents. There is one unit of productive resources. Agent i 's ownership share is denoted r_i and satisfies $r_i \in [0, 1]$, and the productive resources are entirely owned by the agents, that is, $\sum_{i \in \mathcal{N}} r_i = 1$. Each agent i 's consumption utility when allocated $q \in [0, 1]$ units and when its type is θ_i is denoted $V(\theta_i, q)$. The type θ_i is agent i 's private information. It is an independent draw from its type distribution F_i with support $[\underline{\theta}_i, \bar{\theta}_i]$ and density f_i that is positive on the support. We consider two possibilities for the agents' preferences, which we analyze separately. In the case of *constant marginal values*, we have $V(\theta_i, q) = q\theta_i$, and in the case of *decreasing marginal values*, we have $V(\theta_i, q) = q\theta_i - q^2/2$.¹⁶ Type distributions and payoff functions are common knowledge. We assume that V and the supports $[\underline{\theta}_i, \bar{\theta}_i]$ are such that every agent has nonnegative marginal utility for every possible allocation, which means that for every $i \in \mathcal{N}$, $\left. \frac{\partial V(\underline{\theta}_i, q)}{\partial q} \right|_{q=1} \geq 0$. For constant marginal values, this simply means that we require that $\underline{\theta}_i \geq 0$, whereas for decreasing marginal values, it boils down to assuming that $\underline{\theta}_i \geq 1$ for all i . This means that for any feasible q and any $\theta > \hat{\theta}$, we have $V(\theta, q) > V(\hat{\theta}, q)$; that is, consumption utility is increasing

and multi-dimensional types but assume equal weights.

¹⁶See also Choné et al. (2024), who in an extension of their procurement problem, analyze a setting in which the suppliers' consumption utility is quadratic.

in type. Observe that irrespective of whether marginal values are constant or decreasing, $\frac{\partial^2}{\partial \theta \partial q} V(\theta_i, q) = 1$, which ensures that the usual implications of incentive compatibility hold in both cases (see footnote 18).

Agent i 's bargaining (or welfare) weight is denoted $w_i \in [0, 1]$, with at least one agent having a positive weight. For the case with $n = 2$, we let $r \equiv r_1$, so that 2's ownership is $1 - r$, and, normalizing the bargaining weights by dividing by $w_1 + w_2$, we use $w \equiv \frac{w_1}{w_1 + w_2}$ to denote agent 1's weight, with the implication that agent 2's weight is $1 - w$.

2.2 Mechanism

Bargaining among the agents is modeled as being intermediated by a possibly fictitious designer that chooses a bargaining mechanism. Specifically, a direct mechanism is given as $\langle \mathbf{Q}, \mathbf{M} \rangle$, where $\mathbf{Q} : \times_{i \in \mathcal{N}} [\underline{\theta}_i, \bar{\theta}_i] \rightarrow [0, 1]^n$ satisfying $\sum_{i \in \mathcal{N}} Q_i(\boldsymbol{\theta}) = 1$ is the *allocation rule* and $\mathbf{M} : \times_{i \in \mathcal{N}} [\underline{\theta}_i, \bar{\theta}_i] \rightarrow \mathbb{R}^n$ is the *payment rule*. The mechanism is called *direct* because it asks every agent to report its type. Given $\langle \mathbf{Q}, \mathbf{M} \rangle$ and assuming truthful reporting by agents other than agent i , the interim expected allocation and payment of agent i when its report is θ_i are $q_i(\theta_i) \equiv \mathbb{E}_{\boldsymbol{\theta}_{-i}}[Q_i(\boldsymbol{\theta})]$ and $m_i(\theta_i) \equiv \mathbb{E}_{\boldsymbol{\theta}_{-i}}[M_i(\boldsymbol{\theta})]$. When its type is θ_i , its allocation is Q_i and its payment is M_i , agent i 's payoff from the mechanism is $V(\theta_i, Q_i) - M_i$. Given ownership share r_i and type θ_i , the value of i 's outside option is $V(\theta_i, r_i)$.

A prominent allocation rule is the ex post efficient allocation rule, which we denote by $\mathbf{Q}^e(\boldsymbol{\theta})$. This allocation rule maximizes $\sum_{i \in \mathcal{N}} V(\theta_i, Q_i)$. For example, in the case of constant marginal values, $Q_i^e(\boldsymbol{\theta}) = 1$ if and only if $\theta_i > \max \boldsymbol{\theta}_{-i}$, and $Q_i^e(\boldsymbol{\theta}) = 0$ otherwise, where ties have probability zero and can be broken arbitrarily. We denote by $q_i^e(\theta_i) \equiv \mathbb{E}_{\boldsymbol{\theta}_{-i}}[Q_i^e(\boldsymbol{\theta})]$ agent i 's interim expected allocation under the ex post efficient allocation rule.

The mechanism $\langle \mathbf{Q}, \mathbf{M} \rangle$ satisfies (*Bayes Nash*) *incentive compatibility* (IC) if for all $i \in \mathcal{N}$ and all $\theta_i, \hat{\theta}_i \in [\underline{\theta}_i, \bar{\theta}_i]$,¹⁷

$$\mathbb{E}_{\boldsymbol{\theta}_{-i}}[V(\theta_i, Q_i(\boldsymbol{\theta}))] - m_i(\theta_i) \geq \mathbb{E}_{\boldsymbol{\theta}_{-i}}[V(\theta_i, Q_i(\hat{\theta}_i, \boldsymbol{\theta}_{-i}))] - m_i(\hat{\theta}_i).$$

It satisfies interim individual rationality (IR) if for all $i \in \mathcal{N}$ and all $\theta_i \in [\underline{\theta}_i, \bar{\theta}_i]$,

$$u_i(\theta_i) \equiv \mathbb{E}_{\boldsymbol{\theta}_{-i}}[V(\theta_i, Q_i(\boldsymbol{\theta}))] - m_i(\theta_i) - V(\theta_i, r_i) \geq 0 \quad (1)$$

holds, where $u_i(\theta_i)$ is agent i 's interim expected net payoff from participating in the mechanism when its type is θ_i , with “net” meaning net of the outside option $V(\theta_i, r_i)$. An immediate implication of IC is that the interim expected allocation $q_i(\cdot)$ must be nondecreasing.¹⁸

¹⁷As shown in the Online Appendix, for both constant and decreasing marginal values, we have equivalence of Bayesian and dominant strategy incentive compatibility. While this result is known for the case of constant marginal values, it is, to our knowledge, new for the case of decreasing marginal values.

¹⁸To see this, observe that IC for type θ implies that $V(\theta, q(\theta)) - m_i(\theta) \geq V(\theta, q(\hat{\theta})) - m_i(\hat{\theta})$, while IC

By the revelation principle, the focus on direct mechanisms that satisfy IC and IR is without loss of generality. The problem of the designer is to choose $\langle \mathbf{Q}, \mathbf{M} \rangle$ to maximize the weighted sum of the agents' ex ante expected payoffs:

$$\max_{\mathbf{Q}, \mathbf{M}} \mathbb{E}_{\theta} \left[\sum_{i \in \mathcal{N}} w_i (V(\theta_i, Q_i(\theta)) - M_i(\theta)) \right], \quad (2)$$

subject to IC and IR and a no-deficit constraint, which is to say that the designer does not pour any money into the exchange.

The description of the bargaining mechanism is now almost complete. The cases that remain to be addressed are those in which the mechanism designer runs a budget surplus after solving the constrained maximization problem with the objective in (2) when multiple agents have the maximal bargaining weight. For these cases, the mechanism needs to determine how that budget surplus is shared among those agents. To this end, let $\boldsymbol{\eta} = (\eta_i)_{i \in \mathcal{N}}$ with $\eta_i \in [0, 1]$, $\sum_{i \in \mathcal{N}} \eta_i = 1$, and $\eta_i = 0$ if $w_i < \max \mathbf{w}$. Then in case the budget surplus is positive, agent i obtains the share η_i of this surplus. Notice that if $n = 2$ and $w = 1/2$, then $\eta \equiv \eta_1$ can be interpreted as agent 1's generalized Nash bargaining weight when agents 1 and 2 bargain over the budget surplus, with agent 2's weight being $1 - \eta$.¹⁹

It is useful to be able to vary the degree to which the supports of the agents' type distributions overlap or differ. With that in mind, we use the *shifting support* model, which is defined as follows. We set the length of the agents' supports to 1, that is, for each i , we assume $\bar{\theta}_i = 1 + \underline{\theta}_i$, and assume that whenever $\underline{\theta}_i > 0$, i 's distribution F_i is the same on the support $[\underline{\theta}_i, 1 + \underline{\theta}_i]$ as its *primitive* distribution, denoted F_i^P , whose support is $[0, 1]$.²⁰ That is, for $\theta \in [\underline{\theta}_i, 1 + \underline{\theta}_i]$, we have $F_i(\theta) = F_i^P(\theta - \underline{\theta}_i)$. This also means that $f_i(\theta) = f_i^P(\theta - \underline{\theta}_i)$, where f_i^P is the density of the primitive distribution. We denote by $\mu_i \equiv \mathbb{E}_{\theta_i}[\theta_i]$ and $\mu_i^P \equiv \int_0^1 x dF_i^P(x)$ the expected values.

Denote the agents' virtual type functions as

$$\Psi_i^S(\theta) \equiv \theta + \frac{F_i(\theta)}{f_i(\theta)} \quad \text{and} \quad \Psi_i^B(\theta) \equiv \theta - \frac{1 - F_i(\theta)}{f_i(\theta)},$$

where Ψ_i^S is agent i 's *virtual cost* function and Ψ_i^B is agent i 's *virtual value* function.²¹ The

for type $\hat{\theta}$ implies that $V(\hat{\theta}, q(\theta)) - m_i(\theta) \leq V(\hat{\theta}, q(\hat{\theta})) - m_i(\hat{\theta})$. Subtracting the latter inequality from the former yields $V(\theta, q(\theta)) - V(\hat{\theta}, q(\theta)) \geq V(\theta, q(\hat{\theta})) - V(\hat{\theta}, q(\hat{\theta}))$, which, because the cross partials are positive, holds if and only if $q(\cdot)$ is nondecreasing.

¹⁹Letting $S > 0$ be the budget surplus and assuming that the two agents' outside options when they Nash bargain over the division of the budget surplus are 0, the generalized Nash solution maximizes $p^\eta(S - p)^{1-\eta}$ over $p \in [0, S]$, yielding $p^* = \eta S$ as the maximizer.

²⁰One can interpret the shift in support as reflecting a fixed cost per transaction: suppose there is a fixed cost per transaction $\kappa \geq 0$ and the support of agent 2's type distribution is $[\hat{\theta}, 1 + \hat{\theta}]$. Then its effective support is $[\underline{\theta}, 1 + \underline{\theta}]$ with $\underline{\theta} = \hat{\theta} - \kappa$. Under this interpretation, less (more) overlap means a smaller (bigger) fixed cost κ .

²¹As is well known (see, e.g., Mussa and Rosen (1978) or Bulow and Roberts (1989)), with constant

shifting support model then implies that for $\theta \in [\underline{\theta}_i, 1 + \underline{\theta}_i]$, $\Psi_i^S(\theta) = \Psi_i^{S,P}(\theta - \underline{\theta}_i) + \underline{\theta}_i$ and $\Psi_i^B(\theta) = \Psi_i^{B,P}(\theta - \underline{\theta}_i) + \underline{\theta}_i$, where $\Psi_i^{S,P}$ and $\Psi_i^{B,P}$ denote, respectively, the virtual cost and virtual value associated with the primitive distribution.

2.3 Objects of interest

Given the prominent role that the assumption of efficiency plays in bilateral bargaining models with complete information and the prominence of its impossibility in models with incomplete information, a question of particular interest is under which conditions on the primitives will bilateral bargaining be ex post efficient. Therefore, an object of interest is

$$\mathcal{E}_k(\underline{\theta}) \equiv \{(r, w) \in [0, 1]^2 \mid \text{bargaining is ex post efficient}\},$$

which is the set of ownership structures and bargaining weights such that the bargaining mechanism is ex post efficient, where $k \in \{c, d\}$ stands for constant (c) or decreasing (d) marginal values. Letting

$$\mathcal{W}_k(\underline{\theta}) \equiv \{w \mid \exists r \in [0, 1] \text{ s.t. bargaining is ex post efficient}\}$$

and

$$\mathcal{R}_k(\underline{\theta}) \equiv \{r \mid \exists w \in [0, 1] \text{ s.t. bargaining is ex post efficient}\}$$

denote the sets of bargaining weights and ownership structures consistent with the bargaining mechanism being ex post efficient, we have, by construction, $\mathcal{E}_k(\underline{\theta}) = \mathcal{W}_k(\underline{\theta}) \cup \mathcal{R}_k(\underline{\theta})$. Of course, in principle each of these sets could be empty. We will also use

$$\mathcal{W}_k^0(\underline{\theta}) \equiv \{w \mid \exists r \in (0, 1] \text{ s.t. bargaining is ex post efficient}\}$$

to denote the set bargaining weights consistent with ex post efficiency when agent 1 has a positive ownership share, that is, $r > 0$. This set is particularly relevant for values of $\underline{\theta}$ sufficiently large that under ex post efficiency, agent 2 consumes everything no matter what is the type realization because then for $r = 0$, the initial allocation is already ex post efficient.

2.4 Discussion of assumptions

We conclude this section with a discussion of our modeling assumptions. Allowing variability in the extent to which the distributions overlap and, if they do not, the gap between them, is essential insofar as with overlapping supports and constant marginal values, one would never obtain Coasian irrelevance of the ownership structure in light of the impos-

marginal values, the virtual value function $\Psi_i^B(\theta_i)$ captures the marginal revenue associated with agent i while the virtual cost function $\Psi_i^S(\theta_i)$ captures the marginal cost associated with i when its type is θ_i .

sibility theorem of Myerson and Satterthwaite (1983) nor, as we will show, Stigler’s mere redistribution of monopoly gains. In contrast, keeping the distributions on the supports fixed is, while convenient, not essential. Many results generalize beyond this specification. However, the comparative statics results with respect to $\underline{\theta}$ and the interpretation of changes in $\underline{\theta}$ as reflecting productivity differentials are easily and conveniently captured with this shifting-support model. While the assumption of constant marginal values is standard in much of the mechanism design literature, it is at odds with some empirical evidence.²² So extending the mechanism design results and methodology to this setting is motivated both by a desire for theoretical generality and empirical applicability.

As elaborated in Loertscher and Marx (2022), the independent private values model is also essential insofar as it is the only known framework in which a tradeoff between social surplus and rent extraction derives from the primitives, that is, without imposing contractual restrictions. Similarly, the mechanism design (“as-if”) approach has the attractive feature that results are derived from primitives in a systematic way, which contrasts with, e.g., dynamic bargaining games, where there are often a multiplicity of equilibria/beliefs and where results often hinge on features such as the order of moves. The mechanism design approach provides upper bounds on what is achievable in equilibrium and proves tractable. The focus on IR rather than ex post individual rationality reflects an assumption of a strong contracting environment—agents can be forced to participate even if they regret it ex post. In spirit, this is in line with Coase’s assumption of well-defined property rights. It stacks the deck in favor of efficient bargaining outcomes. In a similar vein, the assumption that the designer solves the problem in (2) reflects or extends an assumption that is common in models of complete information bargaining, namely that bargaining is efficient, subject to constraints.

Last, while we assume that agents’ types are their private information—i.e., we study a model with incomplete information—bargaining outcomes mirror those under complete information under conditions that will be made precise below. In particular, when the gap between the supports of agents’ type distributions is sufficiently large, bargaining has the complete-information feature that outcomes are always ex post efficient. Indeed, in such cases, expected outcomes from incomplete information bargaining map directly to those from generalized Nash bargaining. Thus, independent private values models have the ability to capture efficiency effects of ownership and bargaining power while still encompassing complete-information-style outcomes in which bargaining is always ex post efficient.

²²For example, in cap-and-trade schemes with scarcity, there should always be some firms who emit zero emissions if marginal values are constant (and distributions are continuous). But this is contradicted by the data; see Fowle and Perloff (2013) and Liu et al. (2026).

3 The bargaining mechanism

We now characterize the bargaining mechanism. We proceed quickly through the steps that are familiar from typical mechanism design problems and elaborate on steps that are less familiar.

Observe that the Lagrangian associated with the designer's problem (2) is

$$\mathcal{L}(\mathbf{Q}, \mathbf{M}, \rho, \gamma) = \sum_{i \in \mathcal{N}} w_i \mathbb{E}_{\boldsymbol{\theta}} [V(\theta_i, Q_i(\boldsymbol{\theta})) - M_i(\boldsymbol{\theta})] + \rho \sum_{i \in \mathcal{N}} \mathbb{E}_{\boldsymbol{\theta}} [M_i(\boldsymbol{\theta})] + \sum_{i \in \mathcal{N}} \gamma_i u_i(\omega_i),$$

where ρ is the Lagrange multiplier on the no-deficit constraint, $\omega_i \in [\underline{\theta}_i, \bar{\theta}_i]$ is a worst-off type of agent i , that is, $\omega_i \in \arg \min_{\theta_i \in [\underline{\theta}_i, \bar{\theta}_i]} u_i(\theta_i)$, and $\gamma_i \geq 0$ is the multiplier on agent i 's IR constraint. The standard approach in mechanism design is to use the implication of incentive compatibility to express the designer's objective as a function of the allocation rule only rather than of the allocation rule *and* the payment rule. To make headway toward that, we first treat ω_i as given, which allows us to proceed along familiar paths even though ω_i will, in general, be endogenous and interior. In the last step, we then show that one can indeed solve the problem by finding the optimal allocation rule first and then back out the associated payment rule using the payoff equivalence theorem.

The usual arguments, including the envelope theorem (Milgrom and Segal, 2002), imply that under incentive compatibility, $u_i(\cdot)$ is differentiable almost everywhere and satisfies what has become known as the *payoff equivalence theorem*: for all $i \in \mathcal{N}$ and all $\theta, \theta' \in [\underline{\theta}_i, \bar{\theta}_i]$,

$$u_i(\theta) = u_i(\theta') + \int_{\theta'}^{\theta} (q_i(y) - r_i) dy. \quad (3)$$

Thus, up to the constant $u_i(\theta')$, i 's net payoff when of type θ is pinned down by the allocation rule. While the relationship in (3) holds for arbitrary θ and θ' , it is most useful when θ' is a worst-off type of agent i , that is, for $\theta' = \omega_i$. Equating the expression in (3), evaluated at $\theta' = \omega_i$, with the definition of $u_i(\theta)$ in (1) and then solving for $m_i(\theta)$ yields $m_i(\theta_i) = \mathbb{E}_{\boldsymbol{\theta}_{-i}} [V(\theta_i, Q_i(\boldsymbol{\theta}))] - V(\theta_i, r_i) - u_i(\omega_i) - \int_{\omega_i}^{\theta_i} (q_i(y) - r_i) dy$. Using $\mathbb{E}_{\boldsymbol{\theta}} [M_i(\boldsymbol{\theta})] = \mathbb{E}_{\theta_i} [m_i(\theta_i)] = \int_{\underline{\theta}_i}^{\bar{\theta}_i} m_i(\theta_i) dF_i(\theta_i)$ and changing the order of integration in the resulting double integral then yields

$$\mathbb{E}_{\boldsymbol{\theta}} [M_i(\boldsymbol{\theta})] = \mathbb{E}_{\boldsymbol{\theta}} [V(\Psi_i(\theta_i, \omega_i), Q_i(\boldsymbol{\theta})) - V(\Psi_i(\theta_i, \omega_i), r_i)] - u_i(\omega_i), \quad (4)$$

where $\Psi_i(\theta_i, \omega_i)$ is the *overall* virtual type function

$$\Psi_i(\theta_i, \omega_i) \equiv \begin{cases} \Psi_i^S(\theta_i) & \text{if } \theta_i \in [\underline{\theta}_i, \omega_i], \\ \Psi_i^B(\theta_i) & \text{if } \theta_i \in [\omega_i, \bar{\theta}_i]. \end{cases} \quad (5)$$

In standard mechanism design problems, such as optimal sale (respectively procurement)

auctions, it is known a priori that $\underline{\theta}_i$ (respectively $\bar{\theta}_i$) is a worst-off type of agent i because the monotonicity of the allocation rule implied by incentive compatibility means that every interior type enjoys positive information rents. In such problems, one then has $\omega_i = \underline{\theta}_i$ (respectively $\omega_i = \bar{\theta}_i$), implying that $\Psi_i(\theta, \omega_i) = \Psi_i^B(\theta)$ (respectively $\Psi_i(\theta, \omega_i) = \Psi_i^S(\theta)$). In contrast, in our setting, neither $\underline{\theta}_i$ nor $\bar{\theta}_i$ will, in general, be worst-off types of agent i . To see this, observe that the integrand in (3) is negative for y such that $q_i(y) < r_i$ and positive otherwise. Hence, the resulting information rent is not necessarily monotone in θ . This is why the overall virtual type function in (5) is required. Nonetheless, in (4) we now have, as desired, an expression for the expected revenue from agent i that, given ω_i , only depends on the allocation rule.

Noting that $M_i(\boldsymbol{\theta})$ enters the Lagrangian twice, once with weight $-w_i$ and once with weight ρ , straightforward algebra shows that the Lagrangian can be rewritten as

$$\begin{aligned} \hat{\mathcal{L}}(\mathbf{Q}, \rho, \boldsymbol{\gamma}) &= \rho \mathbb{E}_{\boldsymbol{\theta}} \left[\sum_{i \in \mathcal{N}} (V(\Psi_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i), Q_i(\boldsymbol{\theta})) - V(\Psi_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i), r_i)) \right] \\ &\quad + \sum_{i \in \mathcal{N}} (w_i - \rho + \gamma_i) u_i(\omega_i) + \sum_{i \in \mathcal{N}} r_i w_i \mathbb{E}_{\theta_i}[\theta_i], \end{aligned}$$

where

$$\Psi_{i, \alpha}(\theta_i, \omega_i) \equiv \alpha \theta_i + (1 - \alpha) \Psi_i(\theta_i, \omega_i) \quad (6)$$

is i 's *weighted virtual type* function with weight $\alpha = \frac{w_i}{\rho}$ on its true type θ_i . This is the mechanism design analogue to the weighted marginal revenue that is familiar from Ramsey pricing (see, e.g., Wilson, 1993).²³

Observe that, ignoring “irrelevant” agents that have no initial ownership and are never allocated resources, for $\rho < \max \mathbf{w}$ the solution is unbounded because $\hat{\mathcal{L}}$ would be maximized by giving an agent with the maximum weight an infinite amount of money. Consequently, any solution satisfies $\rho \geq \max \mathbf{w}$. Alternatively, if there exist “irrelevant” agents, then $\rho \geq \max\{w_i \mid i \text{ s.t. } r_i > 0 \text{ or } Q_i(\boldsymbol{\theta}) > 0 \text{ for some } \boldsymbol{\theta}\}$ is required.²⁴

If ρ is the solution value of the Lagrange multiplier of the designer's problem, then the allocation rule $\mathbf{Q}$ of the optimal mechanism maximizes $\hat{\mathcal{L}}(\mathbf{Q}, \rho, \boldsymbol{\gamma})$, subject to the monotonicity constraint imposed by incentive compatibility. Although the agents do not necessarily have the same weights in the designer's objective function, this is conceptually the same problem and procedure as in the second-best problem solved by Myerson and Satterthwaite (1983) (who, when so doing, assume increasing virtual type functions).

²³Letting $P(Y)$ denote the inverse demand function for quantity Y , the marginal revenue is $P(Y) + P'(Y)Y$ and the weighted marginal revenue is $\alpha P(Y) + (1 - \alpha)(P(Y) + P'(Y)Y) = P(Y) + (1 - \alpha)P'(Y)Y$.

²⁴An agent i with $r_i = 0$ and $Q_i(\boldsymbol{\theta}) = 0$ for all $\boldsymbol{\theta}$ has $u_i(\omega_i) = 0$ (all types are equally worst-off), and so such an agent essentially drops out of the Lagrangian and, therefore, does not constrain ρ .

There are, however, two additional complications that arise here. First, even with increasing virtual type functions, the overall virtual type is not monotone around ω_i if $\omega_i \in (\underline{\theta}_i, \bar{\theta}_i)$. This can easily happen for $r_i \in (0, 1)$. In this case, pointwise maximization of $\hat{\mathcal{L}}(\mathbf{Q}, \rho, \gamma)$ over $\mathbf{Q}$ may violate the monotonicity constraint imposed by incentive compatibility. This monotonicity requirement is also violated if F_i is such that $\Psi_{i, w_i/\rho}^S$ or $\Psi_{i, w_i/\rho}^B$ is decreasing in the relevant range (a necessary condition for which is that $\Psi_i^S(\theta)$ or $\Psi_i^B(\theta)$ is decreasing in that range). In any of these cases, the solution involves ironing agent i 's weighted virtual type function as in Myerson (1981). For example, for constant marginal values, the resources are allocated to an agent with the highest *ironed* weighted virtual type, which for agent i is denoted by $\bar{\Psi}_{i, \frac{w_i}{\rho}}(\theta, \omega_i)$, with ties broken through randomization that maintains agents' worst-off types.²⁵ We similarly use $\bar{\Psi}_i^S$ and $\bar{\Psi}_i^B$ to denote agent i 's ironed virtual cost and virtual value functions, and similarly for the weighted versions of these.

The second, and deeper, problem is that the interdependence of $\mathbf{Q}$ and $\boldsymbol{\omega} = (\omega_i)_{i \in \mathcal{N}}$ raises the question of the applicability of the standard mechanism design methodology, whereby the objective is first maximized over monotone allocation rules and then the payment rule is derived based on the optimal allocation rule and the payoff equivalence theorem. Fortunately, the answer is affirmative. As observed by Loertscher and Wasser (2019), the optimal mechanism in a partnership model is characterized by a *saddle point* $(\mathbf{Q}^*, \boldsymbol{\omega}^*)$:²⁶ $\mathbf{Q}^*$ is a monotone allocation rule that maximizes $\mathbb{E}_{\boldsymbol{\theta}}[\sum_{i \in \mathcal{N}} V(\Psi_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i^*), Q_i(\boldsymbol{\theta}))]$ and $\boldsymbol{\omega}^*$ is a minimizer of $\mathbb{E}_{\boldsymbol{\theta}}[\sum_{i \in \mathcal{N}} (V(\Psi_{i, \frac{w_i}{\rho}}(\theta_i, x_i), Q_i^*(\boldsymbol{\theta})) - V(\Psi_{i, \frac{w_i}{\rho}}(\theta_i, x_i), r_i))]$ over $\mathbf{x} = (x_i)_{i \in \mathcal{N}}$ with $x_i \in [\underline{\theta}_i, \bar{\theta}_i]$.

To complete the characterization of the bargaining mechanism, we must also satisfy the no-deficit constraint. This is always possible because by choosing $Q_i(\boldsymbol{\theta}) = r_i$ for all i and all $\boldsymbol{\theta}$, the designer obtains revenue of 0. Moreover, in the limit as ρ goes to infinity, the allocation rule approaches that for the mechanism that maximizes the designer's expected revenue. As just observed, the designer's maximized expected revenue must be nonnegative, and it is positive whenever the problem is such that there is a positive measure of types with mutually beneficial trades. By the continuity and monotonicity of the problem, this then guarantees that there is a smallest value of ρ that satisfies the no-deficit constraint.²⁷

Theorem 1. *The allocation rule of the bargaining mechanism pointwise maximizes $\sum_{i \in \mathcal{N}} V(\bar{\Psi}_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i), Q_i(\boldsymbol{\theta}))$, where ω_i is a worst-off type for agent i and ρ is the smallest*

²⁵As shown in the Online Appendix, even with decreasing marginal values there is no benefit from using randomization ex post.

²⁶Loertscher and Wasser (2019) assume equal weights and identical supports, that is, $w_i = w$ and $[\underline{\theta}_i, \bar{\theta}_i] = [0, 1]$ for all $i \in \mathcal{N}$, but these insights extend to heterogeneous weights and supports and to decreasing marginal values.

²⁷See the Online Appendix for a proof that expected revenue under binding IR for agents' worst-off types is continuous and nondecreasing in ρ .

feasible value of the Lagrange multiplier such that the no-deficit constraint is satisfied.

For constant marginal values, Theorem 1 has the following implication:

Proposition 1. *With constant marginal values, the bargaining mechanism assigns the resources to an agent with the maximum ironed weighted virtual type, $\max_{i \in \mathcal{N}} \bar{\Psi}_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i)$, with a tie-breaking rule that ensures that ω_i is a worst-off type for each agent i and where ρ is equal to the smallest feasible value such that the no-deficit constraint is satisfied. For $n = 2$, ties between agents can be broken arbitrary, and if the virtual values and virtual costs are increasing, then ties between agents' ironed weighted virtual types occur with probability zero.*

Proof. The first sentence follows from Theorem 1. For a proof of the second sentence, see the Online Appendix.

The case of constant marginal values described in Proposition 1 has the bang-bang property that, with probability 1, the allocation to agent i is either 0 or 1. This contrasts with the model with decreasing marginal values, where the allocations are, typically, more continuous as we show next. As mentioned, for decreasing marginal values, we assume that for all $i \in \mathcal{N}$, $\theta_i \geq 1$. This ensures that each agent has a nonnegative marginal value no matter what its type realization, even if it is allocated the entire resource, that is, it implies that there is scarcity. Using a water-filling algorithm (see, e.g., Boyd and Vandenberghe, 2004, p. 245), Theorem 1 has the following implication for decreasing marginal values:

Proposition 2. *With decreasing marginal values, the allocation rule of the bargaining mechanism is*

$$Q_i(\boldsymbol{\theta}) = \max\{0, \bar{\Psi}_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i) - \zeta\},$$

with ζ such that $\sum_{i \in \mathcal{N}} Q_i(\boldsymbol{\theta}) = 1$, where ω_i is a worst-off type for each agent i and ρ is equal to the smallest feasible value such that the no-deficit constraint is satisfied. Letting $\mathcal{N}^+$ be the set of agents with a positive allocation, i.e., $Q_i(\boldsymbol{\theta}) > 0$ if and only if $i \in \mathcal{N}^+$, then $\zeta = \frac{1}{|\mathcal{N}^+|} \sum_{j \in \mathcal{N}^+} \bar{\Psi}_{j, \frac{w_j}{\rho}}(\theta_j, \omega_j) - \frac{1}{|\mathcal{N}^+|}$. If $n = 2$, then $Q_1(\boldsymbol{\theta}) = \min\{1, \max\{0, \frac{1}{2}(1 + \bar{\Psi}_{1, \frac{w_1}{\rho}}(\theta_1, \omega_1) - \bar{\Psi}_{2, \frac{w_2}{\rho}}(\theta_2, \omega_2))\}\}$ and $Q_2(\boldsymbol{\theta}) = 1 - Q_1(\boldsymbol{\theta})$.

Interestingly, tie-breaking is not an issue in the decreasing marginal values case. For example, if all agents have the same ironed weighted virtual types, then each agent receives an equal share of the resources. The continuous allocation rule and absence of tie-breaking renders the case of decreasing marginal values in some ways more tractable than the case of constant marginal values.

4 Ex post efficient bilateral bargaining

With the general description of the bargaining mechanism in hand, we now analyze under which conditions bilateral bargaining, that is, bargaining with $n = 2$, is ex post efficient under the assumption of constant marginal values. The objects of interest here are, therefore, $\mathcal{E}_c(\underline{\theta})$, $\mathcal{R}_c(\underline{\theta})$, and $\mathcal{W}_c(\underline{\theta})$.

As mentioned, for the case of bilateral bargaining, we denote ownership by $(r, 1 - r)$ and bargaining weights by $(w, 1 - w)$, and we set $[\underline{\theta}_1, \bar{\theta}_1] = [0, 1]$, which means that F_1 is identical to agent 1's primitive distribution, and $\underline{\theta}_2 = \underline{\theta} \geq 0$ (and $\bar{\theta}_2 = 1 + \underline{\theta}$). We say that bargaining is ex post efficient if given r , w , $\underline{\theta}$, and the agents' type distributions, the allocation rule of the bargaining mechanism is $\mathbf{Q}^e(\boldsymbol{\theta})$. Consequently, for $\underline{\theta} \geq 1$ any trade that occurs under ex post efficiency ships resources from agent 1 to agent 2 no matter what r is. Therefore, for ex post efficiency, agent 1's ownership r determines the size of the asset over which the two agents negotiate and thus effectively corresponds to a rescaling of an asset of size 1 to an asset of size r . Agent 1 is the sole owner and seller of this rescaled asset, with agent 2 being the buyer. Of course, if $r = 0$, there is nothing to be negotiated over. Therefore, for $\underline{\theta} \geq 1$, the ownership structure r has no impact on whether the bargaining mechanism is ex post efficient. As a result, we have $\mathcal{R}_c(\underline{\theta}) = [0, 1]$ for all $\underline{\theta} \geq 1$.

This also means that, for $\underline{\theta} \geq 1$ and $r > 0$, whether bargaining is ex post efficient depends only on the bargaining weights w and $1 - w$. If agent 1 has all the bargaining power, that is, $w = 1$, then by Proposition 1 trade occurs if and only if $\bar{\Psi}_2^B(\theta_2) \geq \theta_1$, which is satisfied for all types if and only if $\bar{\Psi}_2^B(\underline{\theta}) \geq 1$. Conversely, if agent 2 has all the bargaining power, that is, $w = 0$, then trade occurs if and only if $\theta_2 \geq \bar{\Psi}_1^S(\theta_1)$, which holds for all types if and only if $\underline{\theta} \geq \bar{\Psi}_1^S(1)$. As seems intuitive and is formally shown in the proof of the lemma below, extremal bargaining weights impose the tightest conditions for bargaining to be ex post efficient, which means that $\bar{\Psi}_2^B(\underline{\theta}) \geq 1$ and $\underline{\theta} \geq \bar{\Psi}_1^S(1)$ are necessary and sufficient for ex post efficient bargaining for all $w \in [0, 1]$. This allows us to define the threshold value for $\underline{\theta}$ such that bargaining is ex post efficient. Specifically, defining

$$\tau_c^*(0) \equiv \bar{\Psi}_1^S(1), \quad \tau_c^*(1) \equiv 1 - \bar{\Psi}_2^{B,P}(0), \quad \text{and} \quad \tau_c^* \equiv \max\{\tau_c^*(0), \tau_c^*(1)\},$$

we have the following result:

Lemma 1. *With constant marginal values and $\underline{\theta} \geq 1$, bargaining is ex post efficient for all $w \in [0, 1]$ if and only if $\underline{\theta} \geq \tau_c^*$.*

Proof. See Appendix A.

For example, if Ψ_1^S and Ψ_2^B are monotone, then the threshold value for $\underline{\theta}$ defined in Lemma 1 can be written as $\tau_c^* = 1 + \max\left\{\frac{1}{f_1(1)}, \frac{1}{f_2^P(0)}\right\}$. If $\underline{\theta} < \tau_c^*$, then bargaining power matters for

both the size and the distribution of surplus.²⁸ Conversely, and equivalently, we can define threshold values for the agents' bargaining weights. For any $\underline{\theta} \in [1, \tau_c^*(0)]$, let $\underline{w}(\underline{\theta})$ be such that $\bar{\Psi}_{1, \frac{\underline{w}(\underline{\theta})}{1-\underline{w}(\underline{\theta})}}^S(1) = \underline{\theta}$, and otherwise let $\underline{w}(\underline{\theta}) = 0$. And for any $\underline{\theta} \in [1, \tau_c^*(1)]$, let $\bar{w}(\underline{\theta})$ be such that $1 = \bar{\Psi}_{2, \frac{1-\bar{w}(\underline{\theta})}{\bar{w}(\underline{\theta})}}^B(\underline{\theta})$, and otherwise let $\bar{w}(\underline{\theta}) = 1$. Then we have the following result:

Lemma 2. *For $\underline{\theta} \geq 1$ and $r > 0$, bargaining is ex post efficient if and only if $w \in [\underline{w}(\underline{\theta}), \bar{w}(\underline{\theta})]$, where $\underline{w}(\underline{\theta})$ and $\bar{w}(\underline{\theta})$ exist and are unique, with $\underline{w}(1) = \bar{w}(1) = 1/2$, $\underline{w}(\underline{\theta})$ decreasing, and $\bar{w}(\underline{\theta})$ increasing.*

Proof. See Appendix A.

If Ψ_1^S is increasing, then $\underline{w}(\underline{\theta}) = \frac{1+(1-\underline{\theta})f_1(1)}{2+(1-\underline{\theta})f_1(1)}$, and if Ψ_2^B is increasing, then $\bar{w}(\underline{\theta}) = \frac{1}{2+(1-\underline{\theta})f_2^P(0)}$. This implies that, for Ψ_1^S and Ψ_2^B monotone, $\underline{w}(\underline{\theta})$ is concave for $\underline{\theta} \in [1, \tau_c^*(0)]$ and $\bar{w}(\underline{\theta})$ is convex for $\underline{\theta} \in [1, \tau_c^*(1)]$. This is illustrated in Figure 1(c). It assumes that both distributions are uniform, which implies monotone virtual types. For the uniform-uniform case, we have $\tau_c^*(0) = \tau_c^*(1) = 2$, but with other distributions, $\tau_c^*(0) \neq \tau_c^*(1)$ can easily occur.

With this backdrop, we now present our main result, which states that for bilateral bargaining and constant marginal values, depending on the differences in the supports of the two agents' type distributions, the set of ownership and bargaining weights that are consistent with ex post efficiency, $\mathcal{E}_c(\underline{\theta})$, encompasses both non-Coasian relevance of initial ownership as well as Coasian irrelevance, and it encompasses both the Galbraithian importance of countervailing power as well as the mere redistributive effects stipulated by Stigler.

Theorem 2. *In the bilateral bargaining setup with constant marginal values, we have:*

1. *non-Coasian relevance of ownership and a requirement of equal bargaining weights when supports overlap:*

for $\underline{\theta} < 1$, $\mathcal{E}_c(\underline{\theta})$ is defined by a nondegenerate interval of ownership $[\underline{r}(\underline{\theta}), \bar{r}(\underline{\theta})]$ together with equal bargaining weights, where the interval has bounds away from 0 and 1 and includes the ownership that equalizes agents' worst-off types, $r^(\underline{\theta})$, which is decreasing in $\underline{\theta}$, and where the lower and upper bounds of the interval converge to 0 as $\underline{\theta}$ goes to 1, giving us $\mathcal{E}_c(\underline{\theta}) = ([\underline{r}(\underline{\theta}), \bar{r}(\underline{\theta})], 1/2)$;*

2. *Coasian irrelevance of ownership (away from $r = 0$) and a requirement of sufficiently close to equal bargaining weights when supports do not overlap but the gap is sufficiently small:*

²⁸This generalizes to a setup with interior ownership the insight from Loertscher and Marx (2022) that the incomplete information framework has the property that bargaining weights do not only affect the distribution but also the size of expected surplus.

for $\underline{\theta} \in [1, \tau_c^*)$, $\mathcal{E}_c(\underline{\theta})$ is the union of two components, the first of which is $r = 0$ together with any bargaining weights, and the second of which is defined by ownership $r \in (0, 1]$ together with bargaining weights in a nondegenerate interval $[\underline{w}(\underline{\theta}), \bar{w}(\underline{\theta})] \subset [0, 1]$, where the bounds converge to $1/2$ as $\underline{\theta}$ approaches 1 from above, giving us $\mathcal{E}_c(\underline{\theta}) = (0, [0, 1]) \cup ((0, 1], [\underline{w}(\underline{\theta}), \bar{w}(\underline{\theta})])$;

3. Coasian irrelevance of ownership and a merely redistributive role of bargaining weights when the gap in the supports is sufficiently large:

for $\underline{\theta} \geq \tau_c^*$, $\mathcal{E}_c(\underline{\theta}) = ([0, 1], [0, 1])$, which encompasses all ownership and bargaining weights.

Proof. See Appendix A.

As stated in Theorem 2, our bargaining framework with independent private values encompasses Galbraith’s notion of *countervailing power* for $\underline{\theta} < \tau_c^*$, and Stigler’s mere redistribution of monopoly gains otherwise, and it encompasses non-Coasian relevance of the distribution of property rights for $\underline{\theta} < 1$, and otherwise Coasian irrelevance. Part 3 of the theorem addresses the case in which $\underline{\theta} \geq \tau_c^*$, where, as in standard models of complete information bargaining, bargaining is always ex post efficient. In this case, even though agents have independent private values, bargaining is observationally equivalent to generalized Nash bargaining.²⁹ Indeed, as described in footnote 19, when bargaining weights are equal, the sharing parameter η plays the role of agent 1’s generalized Nash bargaining weight. This captures situations that are reflective of the idea expressed, for example, by Stigler (1954, p. 13) that changes in bargaining power “merely redistribute[,] monopoly gains” and provides an instance and a formalization of a setting with sufficiently “little” private information that bargaining outcomes mirror those with complete information. Importantly, though, while the incomplete information bargaining appears observationally equivalent to generalized Nash bargaining when the buyer’s value is $\underline{\theta}_2$ and the seller’s is $\bar{\theta}_1$, the difference $\underline{\theta}_2 - \bar{\theta}_1$ times r is only a lower bound on the actual and expected gains from trade, which are respectively $\theta_2 - \theta_1$ and $\mu_2 - \mu_1$ times r .³⁰ In contrast, parts 1 and 2 of the theorem address cases in which

²⁹The bargaining outcome when $\underline{\theta} \geq \tau_c^*$ has agent 1 sell its r units to agent 2 at a per-unit price of $1 + \eta(\underline{\theta} - 1)$, giving agent 1 an expected net payoff of $r(1 + \eta(\underline{\theta} - 1) - \mu_1)$ and agent 2 an expected net payoff of $r(\mu_2 - (1 + \eta(\underline{\theta} - 1)))$. In contrast, away from ex post efficiency, it is possible for there to be trade from agent 2 to agent 1 with nonoverlapping supports. We return to this in Section 5.

³⁰Another, related qualification regarding the equivalence of complete and incomplete information bargaining for $\underline{\theta} \geq \tau_c^*$ applies when, at an ex ante stage, the agents make noncontractible investments that improve their type distributions. With incomplete information, efficient bargaining implies efficient investment whereas with complete information, hold-up from bargaining induces inefficient investments as in the theory of the firm in the tradition of Grossman and Hart (1986) and Hart and Moore (1990). For formalizations of this point, see, for example, Milgrom (2004), Krämer and Strausz (2007), and Liu et al.

ex post efficiency is not guaranteed. For $\underline{\theta} \leq 1$, $w = 1/2$ is necessary for bargaining to be ex post efficient, and for $\underline{\theta} \in (1, \tau_c^*)$, bargaining weights sufficiently close to being equal are necessary. Parts 1 and 2 of the theorem thus provide a vindication of sorts of Galbraith’s (1954, p. 1) thought that countervailing power is of “substantial, perhaps central, importance” to economics.³¹

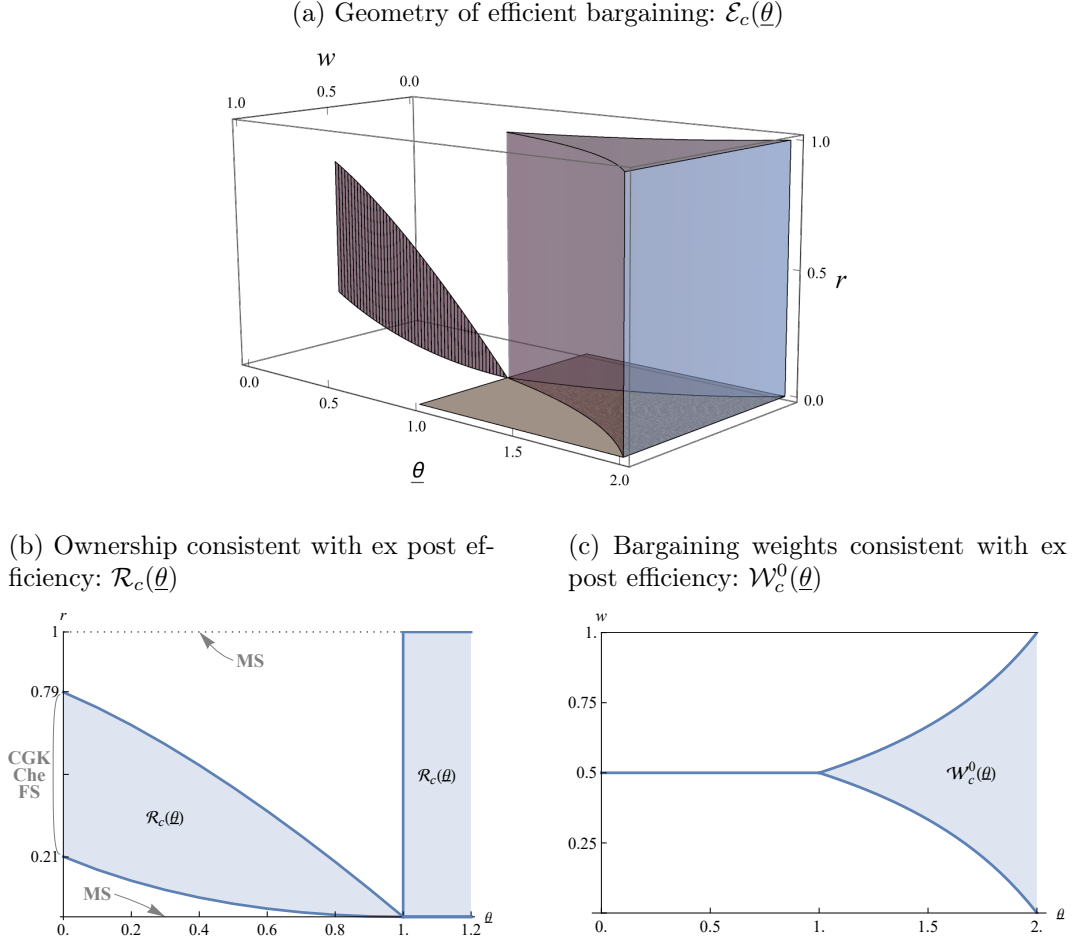


Figure 1: Ownership and bargaining weights consistent with ex post efficiency. Assumes constant marginal values and uniformly distributed types for agent 1 on $[0, 1]$ and for agent 2 on $[\underline{\theta}, 1 + \underline{\theta}]$.

Figure 1 illustrates Theorem 2. Panel (a) shows $\mathcal{E}_c(\underline{\theta})$, panel (b) depicts $\mathcal{R}_c(\underline{\theta})$, and panel

(2026).

³¹The notion of *countervailing power*, introduced by Galbraith (1952), has widespread appeal but has been difficult to conceptualize in models with complete information without restricting the contracting space. It features prominently in antitrust practice. For example, OECD (2011, pp. 50–51) and OECD (2007, pp. 58–59) raise the possibility of a role for collective negotiation and group boycotts in counterbalancing the market power of providers of payment card services. The U.S. DOJ and FTC recognize the potential benefits from allowing physician network joint ventures in their 1996 “Statement of Antitrust Enforcement Policy in Health Care.” Krueger (2018) discusses the benefits to workers of market features that boost worker bargaining power and counterbalance monopsony power. The Australian competition authority “has identified a range of market failures resulting from ... strong bargaining power imbalance and information asymmetry ... which ultimately cause inefficiencies” (ACCC Dairy Inquiry, 2018, p. xii).

(c) depicts $\mathcal{W}_c^0(\underline{\theta})$ as functions of $\underline{\theta}$ for the case in which F_1 and F_2 are uniform. As indicated in panel (b) with the notation “MS,” the result that when supports overlap (i.e., $\underline{\theta} < 1$), ex post efficiency is not possible if $r \in \{0, 1\}$ was shown by Myerson and Satterthwaite (1983); and as indicated with the notation “CGK, Che, FS,” the result that with identical supports (i.e., $\underline{\theta} = 0$), ex post efficiency is possible for an interval of ownerships around the ownership that equalizes agents’ worst-off types was shown by Cramton et al. (1987), Che (2006), and Figueroa and Skreta (2012).

Panel (c) provides an illustration for the uniform-uniform case, where $\Psi_1^S(1) = 2 = 1 - \Psi_2^{B,P}(0)$. If, instead, $\Psi_1^S(1) \neq 1 - \Psi_2^{B,P}(0)$ were the case, then one of the curves emanating from $(1, 1/2)$ would hit the boundary before the other. Specifically, if $\Psi_1^S(1) < 1 - \Psi_2^{B,P}(0)$, then 0 would be an element of $\mathcal{W}_c^0(\underline{\theta})$ for some $\underline{\theta} < \tau_c^*$, while 1 would be in $\mathcal{W}_c^0(\underline{\theta})$ only for $\underline{\theta} \geq \tau_c^*$, and conversely if $\Psi_1^S(1) > 1 - \Psi_2^{B,P}(0)$. As mentioned, the curvature of the bounds of the set $\mathcal{W}_c^0(\underline{\theta})$ shown in panel (c) holds quite generally: away from zero and one, $\underline{w}(\underline{\theta})$ is concave and $\overline{w}(\underline{\theta})$ is convex if the virtual cost and virtual value functions are increasing.

An immediate implication of Theorem 2 is that, for all $\underline{\theta} \geq 0$, $\mathcal{E}_c(\underline{\theta})$ is nonempty. The theorem also allows us to formalize the notions that bargaining weights are more important for ex post efficiency than the ownership structure insofar as bargaining weights matter for a larger set of the parameter space—for all $\underline{\theta} < \tau_c^*$, there are $w \in [0, 1]$ such that bargaining is not ex post efficient, whereas ownerships $r \in [0, 1]$ such that bargaining fails to be ex post efficient only exist for $\underline{\theta} < 1$ —and that, at the same time, getting bargaining weights “right” is easier than the ownership structure insofar as

$$\cap_{\underline{\theta}} \mathcal{R}_c(\underline{\theta}) = \emptyset \quad \text{and} \quad \cap_{\underline{\theta}} \mathcal{W}_c(\underline{\theta}) = \{1/2\},$$

which is to say that there is no ownership structure that is always part of $\mathcal{E}_c(\underline{\theta})$, whereas $w = 1/2$ always is. That $\mathcal{W}_c(\underline{\theta})$ contains $1/2$ for all $\underline{\theta}$ also means that in this bargaining framework, there is no efficiency–equality tradeoff.

A priori, it seems possible that bargaining weights and ownership shares are, somehow, substitutes in the sense that an increase in r can be offset by a commensurate decrease in w to maintain ex post efficiency. However, Theorem 2 shows that this is not the case for ex post efficiency: for $\underline{\theta} \leq 1$, $\mathcal{W}_c(\underline{\theta})$ only contains $w = 1/2$, while for $\underline{\theta} > 1$, whether bargaining is ex post efficient is independent of r , provided r is positive.

The intuition for why $w = 1/2$ is necessary for ex post efficient bargaining if $\underline{\theta} \leq 1$ is simple. Away from equal weights, the bargaining mechanism discriminates against the agent with the smaller weight because, by Proposition 1, the allocation prioritizes agents on the basis of the (ironed) weighted virtual types. For $\underline{\theta} \leq 1$, with unequal bargaining weights, this prioritization differs from prioritizing agents on the basis of their true types. As the

gap between the supports, that is, $\underline{\theta} - 1 > 0$, increases, unequal bargaining weights lead to less and eventually to no discrimination in the allocation rule. That $\mathcal{R}_c(\underline{\theta})$ is a nonempty, convex subset of $(0, 1)$ for $\underline{\theta} \in [0, 1)$ follows from the fact that if r is such that, under ex post efficiency, both agents have the same worst-off types, then revenue under ex post efficiency subject to IR is maximized and positive.³² By the continuity of this revenue function, ex post efficiency is then also possible for a convex set of ownership structures around the revenue-maximizing one. As $\underline{\theta}$ approaches 1 from below, the only way that both agents can have the same worst-off type is if their worst-off types approach 1, which requires that r approaches 0 (in which case revenue under ex post efficiency is simply zero because trade vanishes).

5 Pareto frontiers

We now look at bargaining without restricting attention to ex post efficiency. That is, we study how the agents' expected net payoffs depend on the ownership structure, bargaining weights, and $\underline{\theta}$. We continue to focus on bilateral bargaining with constant marginal values. Hence, the allocation rule is as given by Proposition 1.

Increasing $\underline{\theta}$ reduces the allocative distortions arising from the exertion of bargaining power. As an illustration, assume F_2^P is the uniform distribution and $w = 1$, i.e., agent 1 has all the bargaining power. Then given $\underline{\theta} \geq 0$, we have $\Psi_2^B(\theta_2) = 2\theta_2 - (1 + \underline{\theta})$ and $\Psi_2^S(\theta_2) = 2\theta_2 - \underline{\theta}$, implying that $\Psi_2^{B^{-1}}(x) = \frac{x + \underline{\theta} + 1}{2}$ for $x \in [-1 + \underline{\theta}, 1 + \underline{\theta}]$ and $\Psi_2^{S^{-1}}(x) = \frac{x + \underline{\theta}}{2}$ for $x \in [\underline{\theta}, 2 + \underline{\theta}]$. Because the derivative of these functions with respect to $\underline{\theta}$ is less than 1, the probability that there is trade under the mechanism that is optimal for agent 1 increases for any given θ_1 and any $r \in [0, 1]$. Conversely, and even simpler, the probability that there is trade for any given θ_1 also increases in $\underline{\theta}$ when $w = 0$ because $\Psi_1^B(\theta_1)$ and $\Psi_1^S(\theta_1)$ are independent of $\underline{\theta}$ and the probability that θ_2 exceeds $\Psi_1^B(\theta_1)$ and $\Psi_1^S(\theta_1)$ increases in $\underline{\theta}$.

As illustrated in Figure 2, which assumes that $w = 1$ and $r = 0.7$, for some type realizations there is trade from the higher valuing agent to the lower valuing agent. With overlapping supports, shown in panel (a), for $\theta_1 < 0.3$, there can be trade from agent 1 to agent 2 even though agent 1 has the higher type, and for $\theta_1 > 0.3$, there can be trade from agent 2 to agent 1 even though agent 2 has the higher type. As shown in Figure 2(b), there is also a possibility of trade in the “wrong direction” with nonoverlapping supports: even though agent 2's type always exceeds that of agent 1, for some type realizations, agent 2 sells its 0.3 units to agent 1.

³²This insight is the driving force for why, in Cramton et al. (1987) (who assume identical distributions), the set of ownership structures consistent with ex post efficiency is symmetric around equal ownership. Generalizations of this insight to asymmetric distributions with identical supports were obtained by Che (2006) and Figueroa and Skreta (2012). The proof of Proposition 1 shows that it extends to heterogeneous supports.

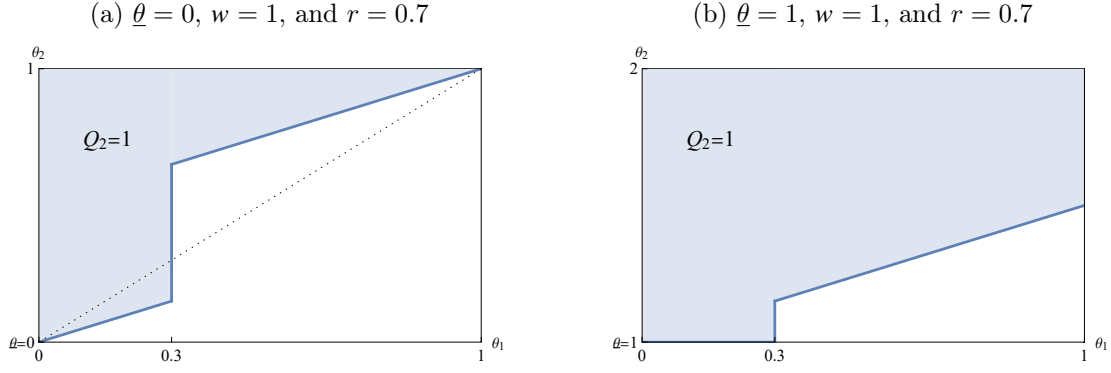


Figure 2: Allocation rule with constant marginal values and $w = 1$. Assumes uniformly distributed types.

Away from extremal bargaining weights, increasing $\underline{\theta}$ has the additional, beneficial effect of making the budget constraint less tight. For example, for $w = 1/2$, $r = 1$, and $\underline{\theta} = 0$, the allocation rule of second-best mechanism is $Q_2 = 1$ if and only if $\theta_2 \geq \underline{\theta} + \theta_1 + 1/4$, while for $\underline{\theta} = 1/4$, it is $Q_2 = 1$ if and only if $\theta_2 \geq \underline{\theta} + \theta_1 + 1/16$.³³ Thus, the strengthening of agent's 2's distribution increases the range of values for agent 2 such that agent 2 is allocated the resource for any given type of agent 1. A comparison of the two panels in Figure 3 illustrates that these comparative statics effects of $\underline{\theta}$ on the second-best allocation rule extend to $r \in (0, 1)$.

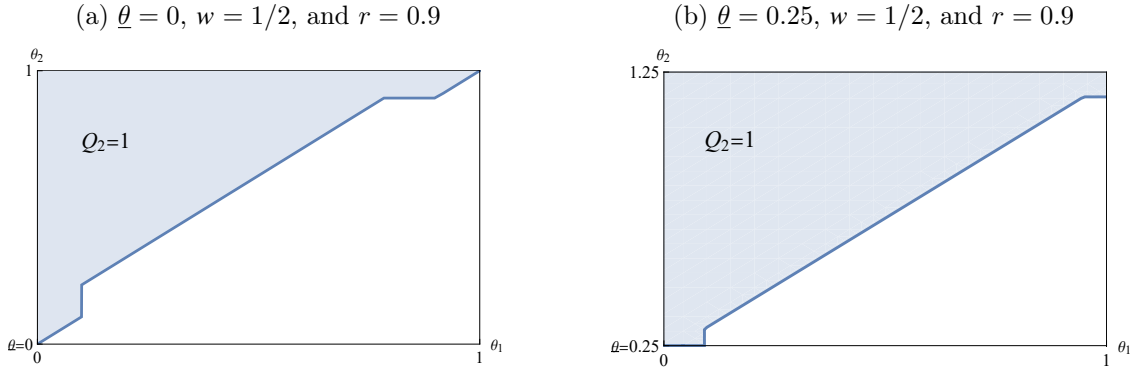


Figure 3: Allocation rule with constant marginal values and equal bargaining weights. Assumes uniformly distributed types.

Let $U_i(r, w) \equiv \mathbb{E}_{\theta_i}[u_i(\theta_i)]$ denote agent i 's expected net payoff given r and w . The frontiers of expected net payoffs, which are illustrated in Figure 4, are the natural objects of interest because they allow us to disentangle the effects of, say, changing r on the performance of the bargaining mechanism from the direct, automatic effects that changes of r have on the agents' utilities via the value of their outside options.

The frontier for a given r is defined by the maximum expected net payoffs that can be

³³In this case, $\rho = 1.5$, $\omega_1 = 1$, and $\omega_2 = 0.25$.

achieved for that r and some $w \in [0, 1]$. With overlapping supports, the frontier point for a given (r, w) is uniquely defined by $(U_1(r, w), U_2(r, w))$, and each bargaining weight w is associated with a unique point on the frontier for a given r , as is illustrated in Figure 4 for the case of uniformly distributed types. As Figure 4 shows, a larger value of agent 1's bargaining weight results in a larger expected net payoff for agent 1 and a smaller expected net payoff for agent 2. But, importantly, the figure also shows the efficiency loss associated with unequal bargaining weights: the farther is w from $1/2$, the smaller is $\sum_{i \in \mathcal{N}} U_i(r, w)$.

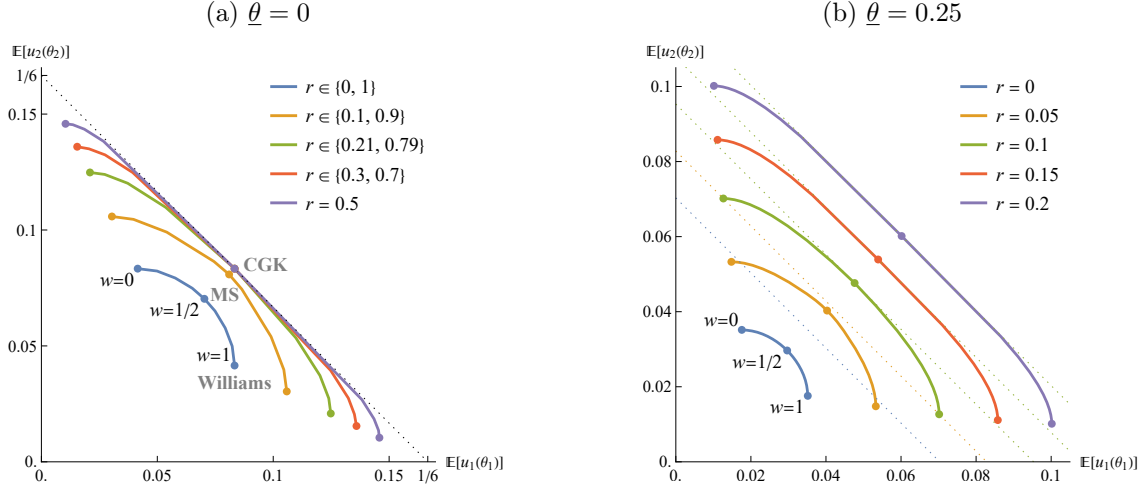


Figure 4: Frontiers for expected net payoffs with constant marginal values. Assumes that agent 1's types are uniformly distributed on $[0, 1]$ and that agent 2's types are uniformly distributed on $[\underline{\theta}, 1 + \underline{\theta}]$, with $\underline{\theta}$ as indicated above each panel. Negatively sloped diagonals reflect expected net payoffs under ex post efficiency, which depends on r in the case of heterogeneous distributions.

To characterize the features of the Pareto frontiers of expected net payoffs as w varies over $[0, 1]$, we denote by $U_{ir}(r, w)$ and $U_{iw}(r, w)$ the derivatives of U_i with respect to r and w , respectively. Then we have:

Proposition 3. *In bilateral bargaining with constant marginal values,*

- (i) *the Pareto frontiers are concave to the origin with slope $\frac{U_{2w}(r, w)}{U_{1w}(r, w)} = -\frac{w}{1-w}$ away from ex post efficiency and slope -1 at ex post efficiency;*
- (ii) *for any given r , the sum of agents' expected net payoffs is maximized when bargaining weights are equal, i.e., $w = 1/2$;*
- (iii) *assuming that Ψ_i^S and Ψ_i^B are increasing, for extremal bargaining weights, as r varies, agents' expected net payoffs move in opposite directions: $U_{1r}(r, 0), U_{2r}(r, 1) > 0$ for r sufficiently close to 1 and $U_{1r}(r, 0), U_{2r}(r, 1) < 0$ for r sufficiently close to 0.*

Proof. See Appendix A for the proofs of (i) and (ii) and Online Appendix B.6 for (iii).

Proposition 3(i) generalizes Loertscher and Marx (2022, Prop. 4) to a partnership setup. It implies that movement toward the equalization of bargaining weights along the expected

net payoff frontier weakly increases social surplus. And, by Proposition 2, for $\underline{\theta} \in [0, 1)$, ex post efficiency is only achieved for full equalization of the bargaining weights.

Proposition 3(iii) has implications for how an agent's expected gain from bargaining depends on its ownership:

Proposition 4. *Assuming that Ψ_i^S and Ψ_i^B are increasing, an agent with no bargaining power expects to gain more from bargaining as ownership approaches an extreme; and, for $F_1 = F_2$, an agent with all the bargaining power expects to gain more from bargaining as ownership moves away from an extreme.*

Proof. The first part follows from Proposition 3(iii), and the proof of the second part is contained in the proof of Proposition 3 in Appendix A.

Figure 4(a) displays the case of identical supports. The case of $r = 0$ and $w = 1/2$, i.e., one seller and one buyer with equal bargaining weights, corresponds to the payoffs associated with the second-best mechanism derived by Myerson and Satterthwaite (1983), which is labeled with “MS” in the figure. Varying w from 0 to 1 while keeping $r = 0$ maps out the frontier for extremal ownership, which was derived by Williams (1987) and is thus labeled “Williams”. Reflecting the impossibility of ex post efficiency with extremal ownership, the entire Williams frontier lies below the ex post efficient frontier, which in Figure 4(a) is given by the line with slope -1 connecting the points $(1/6, 0)$ and $(0, 1/6)$.

Once r increases to 0.21, ex post efficiency becomes possible with equal bargaining weights (labeled “CGK” in the figure). This corresponds to the range $[0.21, 0.79]$ of initial ownership shares for which efficient partnership dissolution is possible in Cramton et al. (1987) when there are two partners with uniformly distributed types on $[0, 1]$. As r increases to 0.5 (by symmetry we need only consider $r \in [0, 0.5]$), the payoff frontier continues to move closer to the ex post efficient frontier, but still only actually touches the frontier for $w = 1/2$. Without the bargaining mechanism from Proposition 1, only the Williams frontier and the points with $w = 1/2$ and $r \in [0.21, 0.79]$ were known.

In Figure 4(b), the supports of the agents' type distributions are only partially overlapping, with $\underline{\theta} = 0.25$. In that case, the sum of the expected net payoffs (or, equivalently, the expected gains from trade) under ex post efficiency depends on the ownership structure because the total expected net payoff varies with r when the means of the agents' type distributions differ. This explains the presence of different ex post efficiency dashed lines in the figure. As the figure shows, ex post efficiency is possible when bargaining weights are equal for r sufficiently close to $1/2$, but ex post efficiency is not possible, even with equal bargaining weights, for $r = 0$ and $r = 0.05$.

A difference arises with nonoverlapping supports because it is then no longer the case that each bargaining weight w is associated with a unique point on the frontier for a given r . As discussed, when $\underline{\theta} > 1$ and $w = 1/2$, bargaining is ex post efficient and akin to generalized Nash bargaining. The weights η and $1 - \eta$ determine the allocation of the expected budget surplus under ex post efficiency when IR binds for the agents' worst-off types. Consequently, when ex post efficiency is possible for bargaining weights other than $w = 1/2$, then the ex post efficient portion of the frontier is defined by two points corresponding to the ex post efficient expected net payoffs for $w < 1/2$ and those for $w > 1/2$, as well as the line segment in between, which represents the expected net payoffs that can be achieved when $w = 1/2$.

Frontiers for nonoverlapping supports are illustrated in Figure 5 for uniformly distributed types. As shown in Figure 5(a) (and Figure 1(c)), for $\underline{\theta} = 1.5$, ex post efficiency is possible for $w \in [1/3, 2/3]$, but not for more extreme values of w , and as shown in Figure 5(b), for $\underline{\theta} = 2$, ex post efficiency is achieved for all $w \in [0, 1]$, in line with Theorem 2.

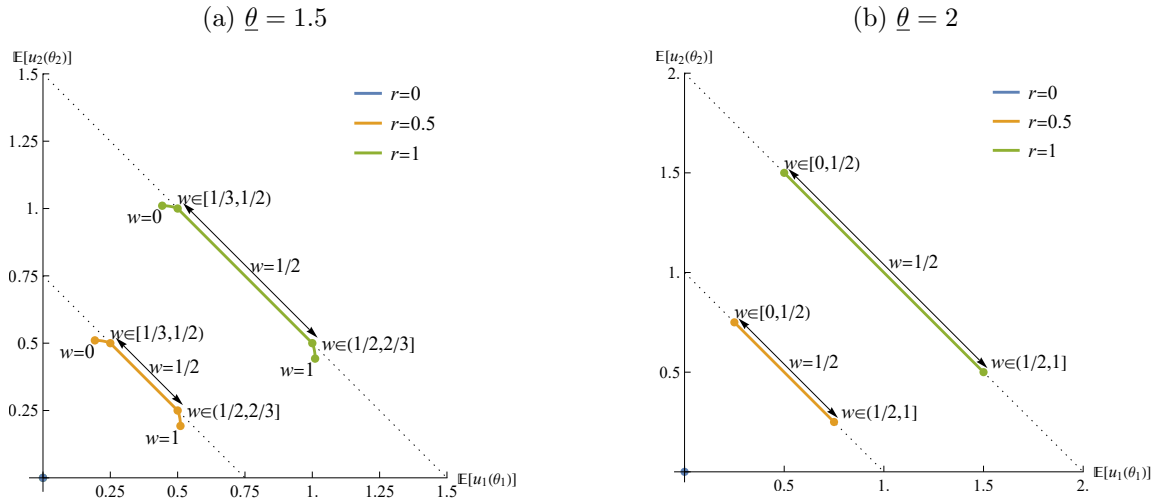


Figure 5: Frontiers for expected net payoffs with constant marginal values. Assumes that agent 1's types are uniformly distributed on $[0, 1]$ and that agent 2's types are uniformly distributed on $[\underline{\theta}, 1 + \underline{\theta}]$, with $\underline{\theta}$ as indicated above each panel. Negatively sloped diagonals reflect expected net payoffs under ex post efficiency, $\mathbb{E}[u_1^e(\theta_1) + u_2^e(\theta_2)]$. When $w = 1/2$, a range of outcomes are possible, as parameterized by $\eta \in [0, 1]$.

Nonlinear effects of the outside option

In Nash bargaining with equal bargaining weights, an agent's payoff is linear in its outside option—with outside options o_i for agents $i = \{1, 2\}$ and total surplus to divide of S , agent i 's Nash bargaining payoff, net of its outside option, is $\frac{1}{2}(S - o_1 - o_2)$. In contrast, with independent private values, this relation is no longer linear even though the expected value of agent i 's outside option, $r_i \mathbb{E}_{\theta_i}[\theta_i]$, is linear in r_i . An agent's outside option affects the agent's worst-off type, which enters both into the IR constraint and into the determination

of the second-best allocation rule. These effects render the relation nonlinear.

As an illustration, Figure 6(a) shows agent 1's expected net payoff as a function of r . As shown in that figure, agents prefer to have higher bargaining power rather than lower bargaining power, all else equal. We can also trace out the frontier for agents' expected net payoffs for a given bargaining weight by varying the resource ownership, as shown in Figure 6(b). That figure highlights that the effects of changes in an agent's outside option vary with the bargaining weights.

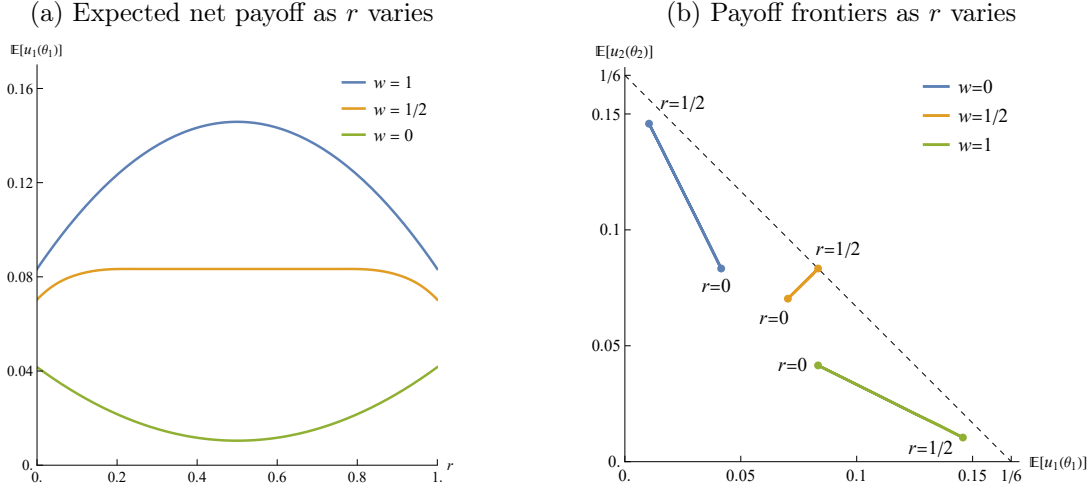


Figure 6: In panel (a), agent 1's expected net payoff as r varies for given w , and in panel (b), frontiers for expected net payoffs as r varies. The negatively sloped diagonal in panel (b) is the ex post efficient frontier. For $w = 1/2$ and $r \in [0.21, 0.79]$, ex post efficiency is achieved. Assumes constant marginal values, $\underline{\theta} = 0$, and that agents' types are uniformly distributed on $[0, 1]$.

Tradeoffs between bargaining power and ownership

As discussed, it may seem intuitive that, to maximize the expected (equally weighted) social surplus of the final allocation, a benevolent social planner may want to offset an agent's bargaining power with less ownership, that is, if an agent has a lot of bargaining power, then the planner may want to give it little ownership, and vice versa. As we have seen, this intuition is not borne out with regard to ex post efficient bargaining because for ex post efficiency, either $w = 1/2$ is required or r does not matter. Away from ex post efficiency, however, both r and w affect the expected social surplus of the final allocation. However, they do so in nonmonotone ways that do not necessarily align with the aforementioned intuition. For example, in the uniform-uniform case with identical supports, when $r = 0.1$, an increase in w from $1/2$ to $2/3$ decreases expected social surplus, but social surplus is more than restored by an increase in r to $1/2$. Thus, in this example an increase in an agent's bargaining power can be offset by an *increase* in this agent's ownership. But this relationship

is not general. For example, for $r = 0.9$, the same increase in w would require a decrease in agent 1's ownership to restore social surplus.

6 Extensions

We now consider extensions and generalizations. Section 6.1 provides an implementation of the bargaining model. In Section 6.2, we analyze multilateral bargaining among more than two agents. In Section 6.3, we return to the case of decreasing marginal values, assuming bilateral bargaining.

6.1 Implementation

We now show that under the assumption of increasing virtual types the bilateral bargaining mechanism can be implemented by a fee-setting mechanism in which a designer charges fees to the agents in exchange for operating a form of Texas shootout: one agent chooses the price and the other chooses whether to sell or buy at that price. Moreover, given known (or estimated) distributions and ownership, observed fee functions are informative regarding w and ρ .

The literature on implementing incomplete information bargaining outcomes has focused on special cases such as ex post efficiency or extremal ownership while imposing parametric restrictions on the distributions. For example, for what corresponds to $w = 1/2$ and $r \in \{0, 1\}$ in our setting, Myerson and Satterthwaite (1983) observe that their second-best mechanism can be implemented using the k -double-auction of Chatterjee and Samuelson (1983) with $k = 1/2$ when both agents draw their types from the uniform distribution with identical supports. Williams (1987) showed that this insight extends insofar as in the uniform-uniform case with $r \in \{0, 1\}$, every point on the Pareto frontier that corresponds to some w can be achieved by the k -double-auction for some value of k .³⁴ Assuming identical distributions and sufficiently symmetric ownership, Cramton et al. (1987) derive an implementation of ex post efficiency using a variant of a k -double-auction applicable to their setup with more than two agents. Given that these results are reasonably well known, we focus in what follows on the new results.

Fee-setting mechanisms

Our implementation of the bilateral bargaining mechanism is an adaptation of the *fee-setting mechanism* of Loertscher and Niedermayer (2019). In our adaption, an intermediary orga-

³⁴Le (2024) shows that for $r_i = 1$ the implementation result obtained by Williams (1987) extends to distributions $F_i(\theta_i) = \theta_i^s$ and $F_j(\theta_j) = 1 - (1 - \theta_j)^b$ with $s, b > 0$ and supports $[0, 1]$.

nizes a variant of a Texas shootout between the agents in exchange for a fee. In line with the assumptions that we make about the designer, the agents' types are not known by the intermediary, but the intermediary knows the distributions from which the agents' types are independently drawn and the agents' bargaining weights. The timing is as follows: First, the intermediary announces (and commits to) fee functions $\phi^B(p)$ and $\phi^S(p)$, which map the price p chosen by agent 1 onto a fee, and fixed payments X_1 and X_2 . Second, agent 1 sets a price p . Third, agent 2 chooses whether to buy agent 1's r units at price p or sell its own $1 - r$ units to agent 2 at price p . Fourth, if agent 2 chooses to buy, then agent 1, who is a seller, pays the fee $r\phi^S(p)$ to the intermediary, and if agent 2 chooses to sell, then agent 1, who is a buyer, pays the fee $(1 - r)\phi^B(p)$ to the intermediary. Last, the intermediary pays fixed amount X_1 to agent 1 and X_2 to agent 2.

As shown in the Online Appendix, this fee-setting game implements the bilateral bargaining mechanism for any $w \in [0, 1]$ and $r \in [0, 1]$.

Proposition 5. *With constant marginal values and monotone virtual costs and virtual values, for any $w \in [0, 1]$, $r \in [0, 1]$, and $\underline{\theta} \geq 0$, there exists a fee-setting mechanism that implements the bilateral bargaining mechanism.*

Proof. See the Online Appendix for a proof by construction.

For intuition, consider the case of $w = 1/2$ and $r \in \mathcal{R}_c(\underline{\theta})$, in which case the bargaining mechanism is ex post efficient. For the fee-setting mechanism, the fee functions need to be chosen to induce agent 1 to set a price equal to θ_1 . Agent 1 with type θ_1 chooses the price p to solve $\max_p \tilde{u}_1(\theta_1, p)$, where $\tilde{u}_1(\theta_1, p)$ is agent 1's interim expected net payoff given price p , ignoring fixed payments:

$$\tilde{u}_1(\theta_1, p) \equiv r(p - \phi^S(p))(1 - F_2(p)) + (\theta_1 - (1 - r)(p + \phi^B(p))) F_2(p) - r\theta_1.$$

The first-order condition for agent 1's maximization problem is

$$\begin{aligned} 0 &= r(1 - \phi^{S'}(p))(1 - F_2(p)) - (1 - r)(1 + \phi^{B'}(p))F_2(p) \\ &\quad + [-r(p - \phi^S(p)) + \theta_1 - (1 - r)(p + \phi^B(p))] f_2(p). \end{aligned}$$

Using the fee functions

$$\phi^S(p) \equiv \int_0^p \frac{r - F_2(x)}{r} dx \quad \text{and} \quad \phi^B(p) \equiv \int_0^p \frac{r - F_2(x)}{1 - r} dx,$$

the first-order condition becomes $(\theta_1 - p) f_2(p) = 0$, which is satisfied, as is the second-order condition, at $p = \theta_1$. Thus, these fees functions induce the ex post efficient allocation.

Agent 1's interim expected net payoff (ignoring fixed payments) is $u_1(\theta_1) \equiv \tilde{u}_1(\theta_1, \theta_1)$, which is minimized at $\hat{\theta}_1 = F_2^{-1}(r)$. Analogously, agent 2's interim expected net payoff

(ignoring fixed payments) is

$$\begin{aligned} u_2(\theta_2) &= \mathbb{E}_{\theta_1}[(1-r)\theta_1 \mid \theta_1 > \theta_2] \Pr(\theta_1 > \theta_2) + \mathbb{E}_{\theta_1}[\theta_2 - r\theta_1 \mid \theta_1 < \theta_2] \Pr(\theta_1 < \theta_2) - (1-r)\theta_2 \\ &= \int_{\theta_2}^1 (1-r)\theta_1 dF_1(\theta_1) + \int_0^{\theta_2} (\theta_2 - r\theta_1) dF_1(\theta_1) - (1-r)\theta_2, \end{aligned}$$

which is minimized at $\hat{\theta}_2 = F_1^{-1}(1-r)$, as required for implementation.

To illustrate, consider uniformly distributed types on $[0, 1]$. Then we have $u_1(\hat{\theta}_1) < 0$ and $u_2(\hat{\theta}_2) > 0$, which means that for individual rationality to be satisfied, agent 1 must be paid a positive fixed fee, while agent 2 can be taxed with a negative fixed fee.

To pin down these fixed fees, note that the intermediary's expected budget surplus under binding IR (not including fixed fees) is

$$\begin{aligned} \Pi^I(r) &= \mathbb{E}_{\theta_1} [r\phi^S(\theta_1)(1 - F_2(\theta_1)) + (1-r)\phi^B(\theta_1)F_2(\theta_1)] + u_1(\hat{\theta}_1) + u_2(\hat{\theta}_2) \\ &= \mathbb{E}_{\theta_1} \left[\int_0^{\theta_1} (r - F_2(x)) dx \right] + u_1(\hat{\theta}_1) + u_2(\hat{\theta}_2). \end{aligned}$$

For the IR and no-deficit to be satisfied, $\Pi^I(r) \geq 0$ is required. This holds for $r \in \mathcal{R}_c(\underline{\theta})$. For example, for uniformly distributed types on $[0, 1]$, we have $\Pi^I(r) = r(1-r) - \frac{1}{6}$. This is nonnegative for $r \in (0.21, 0.79)$. Fixed fees that ensure individual rationality and split the intermediary's expected surplus equally between the agents (consistent with $w = 1/2$) are given by $X_1 = -u_1(\hat{\theta}_1) + \frac{1}{2}\Pi^I(r)$ and $X_2 = -u_2(\hat{\theta}_2) + \frac{1}{2}\Pi^I(r)$, leaving the intermediary with a payoff of zero.

6.2 Multilateral bargaining

In this section, we analyze the ownership structures and bargaining weights that are consistent with ex post efficiency in the model with constant marginal values when there are more than two agents. We let $\mathcal{N} = \mathcal{N}_U \cup \mathcal{N}_D$ consist of agents $i \in \mathcal{N}_U$ and agents $j \in \mathcal{N}_D$. For all $i, j \in \mathcal{N}$ we allow for $F_i \neq F_j$. For all $i \in \mathcal{N}_U$, the support of F_i is $[0, 1]$, while for $j \in \mathcal{N}_D$, F_j has the support $[\underline{\theta}, 1 + \underline{\theta}]$, with $F_j(\theta) = F_j^P(\theta - \underline{\theta})$ for $\theta \in [\underline{\theta}, 1 + \underline{\theta}]$, where F_j^P is j 's primitive distribution with support $[0, 1]$. Let $n_U \equiv |\mathcal{N}_U|$, $n_D \equiv |\mathcal{N}_D|$, and $n = n_U + n_D$.

As we will see, analogs to the results above continue to hold. For example, for ex post efficiency, the agents drawing their types from distributions with the support $[\underline{\theta}, 1 + \underline{\theta}]$ must have equal bargaining weights, and, depending on the supports, bargaining weights of other agents cannot differ by too much. In the interest of space, our focus here is on ex post efficiency, thereby mirroring the analysis in Section 4, even though the analysis away from ex post efficiency extends as well. As in the bilateral setting with constant marginal values, we denote by $\mathcal{R}_c(\underline{\theta})$ the set of ownership structures that are consistent with ex post efficiency.

By Proposition 1, the bargaining mechanism assigns the resource to an agent with the maximum ironed weighted virtual type, $\max_{i \in \mathcal{N}} \bar{\Psi}_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i)$, with ρ equal to the smallest feasible value such that the no-deficit constraint is satisfied. In contrast to the case of $n = 2$, where, as noted in Proposition 1, ties are a zero probability event, for $n > 2$, there can be a positive probability of having more than one agent with the maximum ironed weighted virtual type, so one must address the possibility of ties. In the event of such a tie, the mechanism randomizes over the tied agents with randomization probabilities that are consistent with the agents' worst-off types.

When $n_D = 1$, we let Δ^U denote the set of n_U -dimensional vectors $\mathbf{x}$ such that $(\mathbf{x}, 1 - \sum_{i=1}^{n_U} x_i) \in \Delta$, i.e., $\Delta^U \equiv \{\mathbf{x} \in [0, 1]^{n_U} \mid \sum_{i=1}^{n_U} x_i \leq 1\}$. Further, we let

$$\mathcal{R}_U(\underline{\theta}) \equiv \{\mathbf{r}_U \in \Delta^U \mid (\mathbf{r}_U, 1 - \sum_{i=1}^{n_U} r_{U,i}) \in \mathcal{R}_c(\underline{\theta})\}.$$

With these definitions in hand, the result of Proposition 2 that the set of ownership structures consistent with ex post efficiency converges as $\underline{\theta}$ approaches 1 from below to the singleton set in which the agent with support $[\underline{\theta}, 1 + \underline{\theta}]$ owns all the resources generalizes to a setting with $n_U \geq 2$ agents with support $[0, 1]$ and one agent with support $[\underline{\theta}, 1 + \underline{\theta}]$ by essentially the same logic. And for $\underline{\theta} \geq 1$, again, any ownership structure is consistent with ex post efficiency.

Proposition 6. *With constant marginal values, equal bargaining weights, $n_U \geq 2$, and $n_D = 1$, the set $\mathcal{R}_U(\underline{\theta})$ satisfies: for $\underline{\theta} \in [0, 1)$, $\mathcal{R}_U(\underline{\theta})$ is nonempty with $\mathcal{R}_U \subset [0, 1)^{n_U} \setminus \{\mathbf{0}\}$ and $\lim_{\underline{\theta} \uparrow 1} \mathcal{R}_U(\underline{\theta}) = \{\mathbf{0}\}$; and for $\underline{\theta} \geq 1$, $\mathcal{R}_U(\underline{\theta}) = \Delta^U$.*

Further, we can characterize bargaining weights that are consistent with ex post efficiency. With $n_D \geq 2$, for ex post efficiency to be possible, all agents in $\mathcal{N}_D$ must have the same bargaining weight, and we provide conditions under which the bargaining weights of agents in $\mathcal{N}_U$ are constrained to be equal to or close to those of the agents in $\mathcal{N}_D$. For the purposes of stating Proposition 7, we define $\bar{w}_U \equiv \max_{j \in \mathcal{N}_U \text{ s.t. } r_j > 0} w_j$, which is the maximum bargaining weight among the agents in $\mathcal{N}_U$ that have positive ownership.

Proposition 7. *With constant marginal values, ex post efficiency requires that: (i) all agents in $\mathcal{N}_D$ have the same bargaining weight w_D ; (ii) for $\underline{\theta} \in [0, 1)$, any agent $i \in \mathcal{N}_U$ has $w_i = w_D$; (iii) for $\underline{\theta} \geq 1$, any agent $i \in \mathcal{N}_U$ with $r_i > 0$ has w_i sufficiently close to $\max\{\bar{w}_U, w_D\}$, and with monotone virtual cost and virtual value functions,*

$$\frac{\max\{\bar{w}_U, w_D\} - w_i}{\max\{\bar{w}_U, w_D\}} \leq (\underline{\theta} - 1) f_i(1).$$

Further, for $\underline{\theta} \geq 1$ and monotone virtual cost and virtual value functions, if $n_D = 1$, then bargaining weights are consistent with ex post efficiency if and only if any agent $i \in \mathcal{N}_U$ with

$r_i > 0$ has

$$1 + \left(1 - \frac{w_i}{\max\{w_i, w_D\}}\right) \frac{1}{f_i(1)} \leq \underline{\theta} - \left(1 - \frac{w_D}{\max\{w_i, w_D\}}\right) \frac{1}{f_D(\underline{\theta})}, \quad (7)$$

where we use D as the index for the agent in $\mathcal{N}_D$.

Proof. See the Online Appendix.

As Proposition 7 shows, for overlapping supports, ex post efficiency requires that all agents have the same bargaining weight, consistent with Proposition 2 for the case of only two agents. For nonoverlapping supports, it is still the case that for ex post efficiency all agents with support $[\underline{\theta}, 1 + \underline{\theta}]$ must have the same bargaining weight. But, in that case, the bargaining weights of the agents with support $[0, 1]$ can differ, as long as they do not differ too much from each other and from the common bargaining weight of the agents with support $[\underline{\theta}, 1 + \underline{\theta}]$.

For the case of $n_D = 2$ and $n_U = 1$, Figure 7 provides an illustration of the observation by Makowski and Mezzetti (1993) that for $\underline{\theta} \in (0, 1)$ sufficiently large, ex post efficiency is possible even if the agent in $\mathcal{N}_U$ owns all the resources: the top vertex in the triangle is included in the set of ownership structures consistent with ex post efficiency for $\underline{\theta} = 0.8$ in panel (a) and for $\underline{\theta}$ approaching 1 in panel (b).

(a) Ownership consistent with ex post efficiency: $\mathcal{R}_c(\underline{\theta})$

(b) Ownership consistent with ex post efficiency: $\mathcal{R}_c(\underline{\theta})$ (with limiting case)

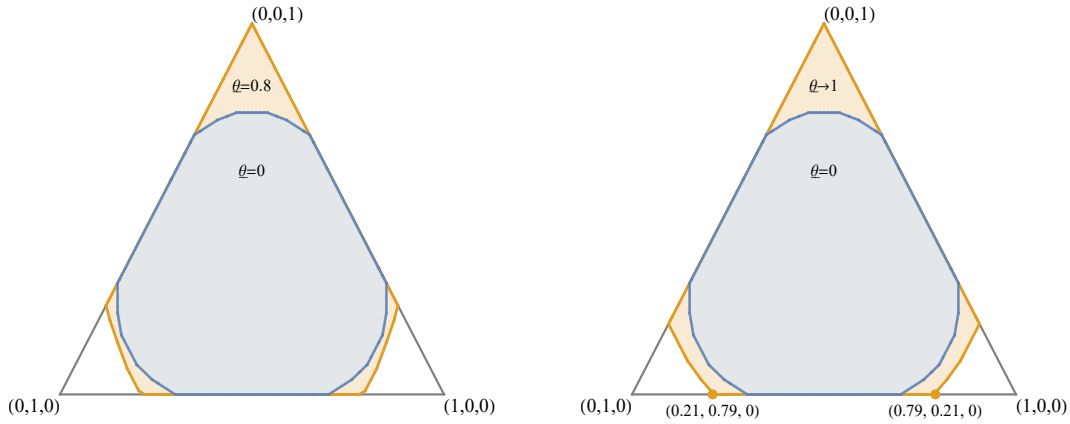


Figure 7: Ownership consistent with ex post efficiency when agents 1 and 2 are in $\mathcal{N}_D$ and agent 3 is in $\mathcal{N}_U$. Assumes constant marginal values and uniformly distributed types on the respective supports.

Figure 7(a) shows that a shift of ownership from an agent in $\mathcal{N}_U$ (agent 3) to an agent in $\mathcal{N}_D$ (agent 1 or 2) can cause ex post efficiency to no longer be possible. For example, starting from the orange boundary in the lower right corner in Figure 7(a), shifting ownership to agent 1 would cause ex post efficiency to no longer be possible (moves the ownership structure into

the white corner triangle). Thus, it may be advantageous to have some ownership by a firm with a weaker distribution to balance the stronger distributions of the agents in $\mathcal{N}_D$.

Related to Figure 7(b), note that when $\underline{\theta} \geq 1$ and $r_3 = 0$, the problem essentially reduces to bilateral bargaining between two symmetric agents, that is, agents 1 and 2. Thus, from Cramton et al. (1987), we know that ex post efficiency is possible if and only if the agents' ownership is sufficiently symmetric. This is reflected in the figure by the fact that the intersection of the orange region labeled " $\underline{\theta} \rightarrow 1$ " with the bottom edge of the triangle, where agent 3's share is zero, spans $(0.21, 0.79, 0)$ to $(0.79, 0.21, 0)$, but excludes more extreme ownership.

6.3 Decreasing marginal values

We now return to bilateral bargaining and to the case of decreasing marginal values. As mentioned, with decreasing marginal values, we let $[\underline{\theta}_1, \bar{\theta}_1] = [1, 2]$ and $\underline{\theta}_2 = \underline{\theta} \geq 1$ (and, again, $\bar{\theta}_2 = 1 + \underline{\theta}$).

Ex post efficiency with decreasing marginal values

With decreasing marginal values, the ex post efficient allocation rule $\mathbf{Q}^e(\boldsymbol{\theta})$ maximizes $\sum_{i \in \mathcal{N}} (Q_i \theta_i - \frac{1}{2} Q_i^2)$. Given the adjusted supports of $[1, 2]$ for agent 1's distribution and $[\underline{\theta}, 1 + \underline{\theta}]$ with $\underline{\theta} \geq 1$ for agent 2, this implies that

$$Q_1^e(\boldsymbol{\theta}) = \max \left\{ 0, \frac{1 + \theta_1 - \theta_2}{2} \right\}$$

and $Q_2^e(\boldsymbol{\theta}) = 1 - Q_1^e(\boldsymbol{\theta})$. The expected budget surplus under binding IR for agents' worst-off types and ex post efficiency is

$$\Pi_d(r) \equiv \sum_{i \in \mathcal{N}} \mathbb{E}_{\boldsymbol{\theta}} \left[\Psi_i(\theta_i, \omega_i) Q_i^e(\boldsymbol{\theta}) - \frac{Q_i^e(\boldsymbol{\theta})^2}{2} \right] - \omega_1 r - \omega_2 (1 - r) + r^2/2 + (1 - r)^2/2,$$

where ω_i is determined as follows: if $r_i \leq q_i^e$, then $\omega_i = \underline{\theta}_i$; if $q_i^e(\bar{\theta}_i) \leq r_i$, then $\omega_i = \bar{\theta}_i$; and otherwise $q_i^e(\omega_i) = r_i$. The effects of this difference are apparent in Figure 8(b), where, in contrast to $\mathcal{R}_c(\underline{\theta})$ in Figure 1(b), $\mathcal{R}_d(\underline{\theta})$ is not convex for all $\underline{\theta}$.

As we saw, for the case of constant marginal values, for $\underline{\theta} \geq 1$, $Q_1^e(\boldsymbol{\theta}) = 0$ for all $\boldsymbol{\theta}$. In contrast, with decreasing marginal values, this threshold is no longer equal to 1. To capture this, we define τ_k^e with $k \in \{c, d\}$ to be the threshold for constant (c) and decreasing (d) marginal values such that for $\underline{\theta} \geq \tau_k^e$, we have $Q_2^e(\boldsymbol{\theta}) = 1$ for all $\boldsymbol{\theta}$. Thus, as shown above $\tau_c^e = 1$, and as we now show, $\tau_d^e = 3$: because $Q_1^e(\boldsymbol{\theta}) = \max\{0, \frac{1 + \theta_1 - \theta_2}{2}\}$, $Q_1^e(\boldsymbol{\theta}) = 0$ for all $\boldsymbol{\theta}$

is equivalent to $\theta_2 \geq 1 + \theta_1$ for all θ , which is equivalent to³⁵

$$\underline{\theta} \geq \tau_d^e = 3.$$

Consequently, if $\underline{\theta} \geq \tau_d^e$, then the ownership structure does not affect whether bargaining is ex post efficient, but only how much is traded—agent 1’s share r —under ex post efficiency.

Theorem 3 below, which is the counterpoint to Theorem 2, is of interest both for its similarities and its differences relative to Theorem 2.

Theorem 3. *In the bilateral bargaining setup with decreasing marginal values, we have:*

1. *non-Coasian relevance of ownership and a requirement of equal bargaining weights when the gap in supports is sufficiently small:*

for $\underline{\theta} = 1$, $\mathcal{E}_d(1) = (\frac{1+\mu_1^P-\mu_2^P}{2}, 1/2)$, and for $\underline{\theta} \in (1, \tau_d^e)$,

$$\mathcal{E}_d(\underline{\theta}) = ((\underline{r}(\underline{\theta}), \bar{r}(\underline{\theta})), 1/2) \cup ((r^{top}(\underline{\theta}), 1), 1/2),$$

where (i) there exists $\hat{\tau} > \min\{\mu_1^P + 1, 2 - \mu_2^P\}$ such that for all $\underline{\theta} \in (1, \hat{\tau})$, there exists decreasing $r^*(\underline{\theta})$ with $0 < \underline{r}(\underline{\theta}) < r^*(\underline{\theta}) < \bar{r}(\underline{\theta}) \leq 1$; and (ii) there exists $\bar{\tau} \in (1, 2)$ such that there exists continuous $r^{top}(\underline{\theta})$ satisfying $r^{top}(\bar{\tau}) = 1$ and for all $\underline{\theta} \in (\bar{\tau}, \tau_d^e)$, $r^{top}(\underline{\theta}) \in (0, 1)$;

2. *Coasian irrelevance of ownership (away from $r = 0$) and a requirement of sufficiently close to equal bargaining weights when the gap is moderate: the result is as in Theorem 2(2) but with bounds $\tau_d^* > \tau_d^e$ and $\underline{\theta} \in [\tau_d^e, \tau_d^*]$;*
3. *Coasian irrelevance of ownership and merely redistributive role of bargaining weights when the gap is sufficiently large: the result is as in Theorem 2(3) but for $\underline{\theta} \geq \tau_d^*$.*

Proof. See Appendix B.8.

As Theorem 3 shows, key characteristics of ownership, bargaining weights, and supports that are consistent with ex post efficiency are preserved with decreasing marginal values, but there are differences. As a point of commonality, for all $\underline{\theta} \geq \tau_d^e$, $\mathcal{R}_d(\underline{\theta}) = [0, 1]$. In a notable difference, as discussed above, the set $\mathcal{E}_d(\underline{\theta})$ is no longer necessarily convex for all $\underline{\theta}$. Theorem 3 establishes that $\mathcal{E}_d(\underline{\theta})$ is nonempty for $\underline{\theta} \in (1, \hat{\tau}] \cup [\bar{\tau}, \infty)$, but leaves the possibility that it is empty for some $\underline{\theta}$, although it is nonempty for all $\underline{\theta}$ in the uniform-uniform case, as illustrated in Figure 8.

We further detail the effects of bargaining power for the case of decreasing marginal values and compare those to the case of constant marginal values. With decreasing marginal

³⁵In contrast to the constant marginal values model, with decreasing marginal values, no overlap in the marginal values is different from no overlap in the type distributions, so τ_d^e differs from the upper bound of support of agent 1’s type distribution.

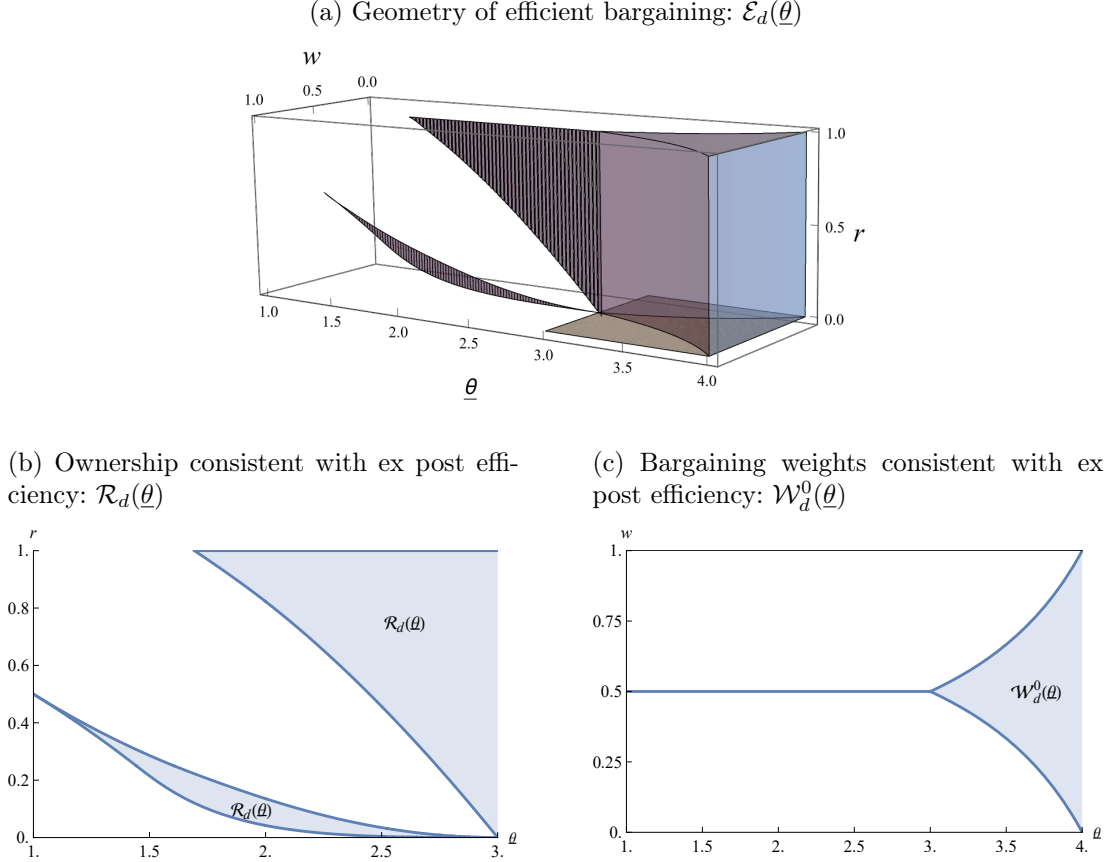


Figure 8: Ownership and bargaining weights consistent with ex post efficiency. Assumes decreasing marginal values and uniformly distributed types for agent 1 on $[1, 2]$ and for agent 2 on $[\underline{\theta}, 1 + \underline{\theta}]$ with $\underline{\theta} \geq 1$.

values, ex post efficiency requires that for $\underline{\theta} \geq \tau_d^e$, $Q_1(\boldsymbol{\theta}) = 0$ for all $\boldsymbol{\theta}$, which requires that $\frac{1}{2}(1 + \bar{\Psi}_{1, \max\{w, 1-w\}}^S(\theta_1) - \bar{\Psi}_{2, \max\{w, 1-w\}}^B(\theta_2)) \leq 0$ for all $\boldsymbol{\theta}$.³⁶ Consequently, for $w = 1$, $1 + \theta_1 - \bar{\Psi}_2^B(\theta_2) \leq 0$ has to hold for all $\boldsymbol{\theta}$. This means that $\bar{\Psi}_2^B(\underline{\theta}) \geq 3$ has to hold. The threshold value for $\underline{\theta}$, denoted by $\tau_d^*(1)$, is then defined by $\bar{\Psi}_2^B(\tau_d^*(1)) = 3$, which means that $\tau_d^*(1) = 3 - \bar{\Psi}_2^{B,P}(0)$. Recalling that for decreasing marginal values, we have $\tau_d^e = 3$, it follows that

$$\tau_d^*(1) = \tau_d^e - \bar{\Psi}_2^{B,P}(0),$$

analogous to the case of constant marginal values. Similarly, when agent 2 has all the bargaining power, that is, $w = 0$, ex post efficiency requires $1 + \bar{\Psi}_1^S(\theta_1) - \theta_2 \leq 0$ for all $\boldsymbol{\theta}$. Thus, ex post efficiency requires $\underline{\theta} \geq 1 + \bar{\Psi}_1^S(2)$, giving us the threshold $\tau_d^*(0) \equiv 1 + \bar{\Psi}_1^S(2) = 2 + \bar{\Psi}_1^{S,P}(1)$. It follows that

$$\tau_d^*(0) = \tau_d^e - 1 + \bar{\Psi}_1^{S,P}(1),$$

³⁶Because agent 1 is always a seller, its worst-off type is $\omega_1 = 2$; and because agent 2 is always a buyer, its worst-off type is $\omega_2 = \underline{\theta}$. Thus, $\Psi_{1,\alpha}(\theta_1, \omega_1) = \Psi_{1,\alpha}^S(\theta_1)$ and $\Psi_{2,\alpha}(\theta_2, \omega_2) = \Psi_{2,\alpha}^B(\theta_2)$.

which is again analogous to the case of constant marginal values. Further, as before, the case of $w \in (0, 1)$ is intermediate, so the threshold $\tau_d^* \equiv \max\{\tau_d^*(0), \tau_d^*(1)\}$ is sufficient. As in the case of constant marginal values, if the virtual cost and virtual value are monotone, then $\tau_d^* = \tau_d^e + \max\{\frac{1}{f_1^P(1)}, \frac{1}{f_2^P(0)}\}$.

We summarize these results in Proposition 8, which highlights commonalities between the setups with constant versus decreasing marginal values.

Proposition 8. *For $k \in \{c, d\}$ for constant and decreasing marginal values, respectively,*

- (i) *for $\underline{\theta} \in [\underline{\theta}_1, \tau_k^e)$, bargaining is ex post efficient if and only if $w = 1/2$ and r is such that $\Pi_k \geq 0$;*
- (ii) *for $\underline{\theta} \in [\tau_k^e, \tau_k^*)$, bargaining is ex post efficient if and only if either $r = 0$ or w satisfies*

$$\overline{\Psi}_{1, \frac{w}{1-w}}^{S,P}(1) \leq 1 + \underline{\theta} - \tau_k^e \quad \text{and} \quad \overline{\Psi}_{2, \frac{1-w}{w}}^{B,P}(0) \geq \tau_k^e - \underline{\theta};$$

- (iii) *for $\underline{\theta} \geq \tau_k^*$, bargaining is ex post efficient for all r and w .*

When virtual costs are not ironed “at the top” and virtual values are not ironed “at the bottom,” we can rewrite the expression in case (ii) of Proposition 8 as follows:

Corollary 1. *For $k \in \{c, d\}$ for constant and decreasing marginal values, respectively, if $\overline{\Psi}_1^{S,P}(1) = \Psi_1^{S,P}(1)$ and $\overline{\Psi}_2^{B,P}(0) = \Psi_2^{B,P}(0)$, then for $\underline{\theta} \in [\tau_k^e, \tau_k^*)$, bargaining is ex post efficient if and only if either $r = 0$ or w satisfies*

$$\frac{1 + (\tau_k^e - \underline{\theta})f_1^P(1)}{2 + (\tau_k^e - \underline{\theta})f_1^P(1)} \leq w \leq \frac{1}{2 + (\tau_k^e - \underline{\theta})f_2^P(0)}.$$

As we have shown, in the decreasing marginal values setup with identical supports, ex post efficiency is possible if and only if $w = 1/2$ and $r = \frac{1+\mu_1-\mu_2}{2}$, which implies that each firm’s ownership must be equal to its expected allocation. This contrasts with the open set of ownerships that allow ex post efficiency with constant marginal values. In either setup, for $\underline{\theta} \geq \tau_k^e$, no restrictions on r are required for ex post efficiency, although there continue to be restrictions on w for ex post efficiency for $\underline{\theta} \in [\tau_k^e, \tau_k^*)$, where τ_k^* is distribution dependent.

Pareto frontiers with decreasing marginal values

Figure 9 illustrates the frontiers for the agents’ expected net payoffs (expected gains from trade) for the setup with decreasing marginal values. Interestingly, the expression for expected net payoffs in the decreasing marginal values setup matches that for the constant marginal values setup, although the quantities and worst-off types differ. To see this, note that under decreasing marginal values, agent i ’s expected payment is

$$\mathbb{E}_{\theta_i}[m_i(\theta_i)] = \mathbb{E}_{\theta}[\Psi_i(\theta_i, \omega_i)Q_1(\theta) - Q_i(\theta)^2/2] - \omega_i r_i + r_i^2/2 - u_i(\omega_i).$$

Thus, agent i 's expected net payoff is

$$\begin{aligned}\mathbb{E}_{\theta_i}[u_i(\theta_i)] &= \mathbb{E}_{\theta}[\theta_i Q_i(\theta) - Q_i(\theta)^2/2 - m_i(\theta_i) - \theta_i r_i + r_i^2/2] \\ &= \mathbb{E}_{\theta}[(\theta_i - \Psi_i(\theta_i, \omega_i))Q_i(\theta) + r_i(\omega_i - \theta_i)] + u_i(\omega_i).\end{aligned}$$

Thus, the quadratic terms drop out, and we have an expression that matches that for the case of constant marginal values, but with allocations and worst-off types that reflect decreasing marginal values.

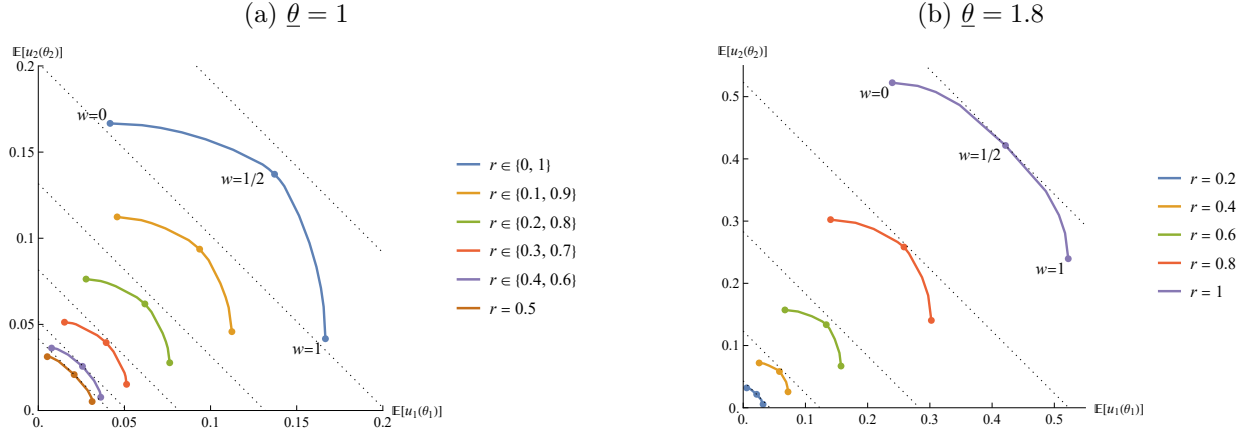


Figure 9: Frontiers for expected net payoffs with decreasing marginal values. Assumes that agent 1's types are uniformly distributed on $[1, 2]$ and that agent 2's types are uniformly distributed on $[\underline{\theta}, 1 + \underline{\theta}]$, with $\underline{\theta}$ as indicated. Negatively sloped diagonals reflect expected net payoffs under ex post efficiency. (For $\underline{\theta} = 1$, ex post efficiency is achieved with $w = 1/2$ and $r = 1/2$, and for $\underline{\theta} = 1.8$, with $w = 1/2$ and $r \in [0.008, 0.192] \cup [0.941, 1]$.)

The similarities and differences between Figure 9(a) and the corresponding figure for constant marginal values, Figure 4(a), highlight features of the bargaining mechanism. Both figures assume identical supports and identical distributions, in which case the expected net payoffs under ex post efficiency (dotted lines in the figures) do not vary with r under constant marginal values, but do so under decreasing marginal values. This is because with constant marginal values, the agents' expected outside options sum to $\mathbb{E}_{\theta}[\theta_1 r + \theta_2(1 - r)]$, which does not vary with r when the agents draw their types from the same distribution. In contrast, under decreasing marginal values, the expected outside option is $\mathbb{E}_{\theta}[\theta_1 r + \theta_2(1 - r) - r^2/2 - (1 - r)^2/2]$, which is strictly concave in r and, with identical distributions, maximized at $r = 1/2$. As a result, the achievable expected net payoffs in Figure 9(a) vary with r and are lowest for $r = 1/2$ and highest for extremal r .

Another key feature of the figures is how far each frontier is below its corresponding dotted ex post efficiency line. In both Figure 4(a) and Figure 9(a), frontiers for more extremal ownership are farther from ex post efficiency than frontiers for more symmetric ownership. With constant marginal values, the frontiers achieve ex post efficiency when $w = 1/2$ and

ownership is sufficiently close to 50–50, whereas with decreasing marginal values, the frontiers only achieve ex post efficiency when ownership is exactly 50–50. This reflects the effects of the IR constraint, which requires that agents’ worst-off types have nonnegative expected payoffs net of their outside option. The sum of the outside options for the agents’ worst-off types is $r\omega_1 + (1-r)\omega_2$ under constant marginal values and $r\omega_1 + (1-r)\omega_2 - r^2/2 - (1-r)^2/2$ under decreasing marginal values. In both cases, these outside options are maximized with extremal worst-off types, with the implication that the no-deficit constraint is “harder” to satisfy when ownership is extremal than when it is more symmetric. The combined result of these two effects is that in Figure 4(a), the frontier for extremal r is the “innermost” frontier, whereas in Figure 9(a), the frontier for extremal r is the “outermost” frontier.

Both Figure 4(b) and Figure 9(b) have asymmetry between the supports, which introduces an additional effect that higher r means higher ownership for the weaker agent, which reduces the expected value of the agents’ outside options, contributing to the shifts in the expected net payoffs under ex post efficiency in those figures.

7 Conclusions

We analyze a model of bargaining with independent private values in which information is always incomplete. We provide conditions under which complete and incomplete information bargaining are observationally equivalent because the bargaining mechanism mimics generalized Nash bargaining. If these conditions are met, then bargaining is always efficient, that is, neither bargaining power nor ownership affects whether the outcome is ex post efficient. This is a formalization and conceptualization of the notion of “little private information.” While ownership structures can affect whether bargaining is efficient, they are, loosely speaking, less important than bargaining power insofar as, typically, there are many ownership structures that are consistent with ex post efficiency, whereas with overlapping supports, ex post efficiency is only possible with equal bargaining power.

The key insights extend to a model with decreasing marginal values in the form of quadratic utility. There are, however, subtle but important difference to the widely studied constant marginal values model as well. As a case in point, with identical supports, there is only one ownership structure that is consistent with ex post efficiency. In addition, bilateral bargaining can be ex post efficient with overlapping supports and extremal ownership.

A Appendix: Proofs

Proof of Lemma 1. The condition for ex post efficiency when $w = 1$, $\bar{\Psi}_2^B(\underline{\theta}) \geq 1$, can be written as $\bar{\Psi}_2^{B,P}(0) + \underline{\theta} \geq 1$, which, noting that $\tau_c^e = 1$, is satisfied for $\underline{\theta} \geq \tau_c^*(1)$. The condition for ex post efficiency when $w = 0$, $\underline{\theta} \geq \bar{\Psi}_1^S(1)$, is satisfied for $\underline{\theta} \geq \tau_c^*(0)$. Turning to the case of $w \in (0, 1)$, a necessary and sufficient condition for bargaining to be ex post efficient, is

$$\bar{\Psi}_{1, \max\{w, 1-w\}}^S(1) \leq \bar{\Psi}_{2, \max\{w, 1-w\}}^B(\underline{\theta}). \quad (\text{A.1})$$

For $w \leq 1/2$, this amounts to $\bar{\Psi}_{1, \frac{w}{1-w}}^S(1) \leq \underline{\theta}$, which holds for $\underline{\theta} \geq \tau_c^*(0)$ because $\bar{\Psi}_{1, \frac{w}{1-w}}^S(\underline{\theta}) \leq \bar{\Psi}_1^S(\underline{\theta})$. For $w > 1/2$, (A.1) is equivalent to $1 \leq \bar{\Psi}_{2, \frac{1-w}{w}}^B(\underline{\theta})$, which holds for $\underline{\theta} \geq \tau_c^*(1)$ because $\bar{\Psi}_{2, \frac{1-w}{w}}^B(\underline{\theta}) \geq \bar{\Psi}_2^B(\underline{\theta})$. Thus, $\underline{\theta} \geq \tau_c^*$ is necessary and sufficient for bargaining to be ex post efficient for all $w \in [0, 1]$. ■

Proof of Lemma 2. First note that $\bar{\Psi}_{1,\alpha}^S(1)$ is decreasing and $\bar{\Psi}_{2,\alpha}^B(\underline{\theta})$ is increasing in α . To see this, note that $\bar{\Psi}_{1,\alpha}^S(1) = \alpha + (1 - \alpha)\bar{\Psi}_1^S(1)$, which is decreasing in α if and only if $\bar{\Psi}_1^S(1) > 1$, which holds because (i) $\bar{\Psi}_1^S(1) > 1$ and (ii) if $\bar{\Psi}_i^S(\theta_1)$ is ironed at $\theta_1 = 1$, then there exists $a \in (0, 1)$ such that $\bar{\Psi}_1^S(1) = \frac{\int_a^1 \Psi_1^S(x) dF_1(x)}{1 - F_1(a)} = \frac{1 - aF_1(a)}{1 - F_1(a)} > 1$. Analogously, one can show that $\bar{\Psi}_{2,\alpha}^B(\underline{\theta})$ is increasing in α using that if $\bar{\Psi}_2^B$ is ironed at $\underline{\theta}$, then there exists $a \in (\underline{\theta}, 1 + \underline{\theta})$ such that $\bar{\Psi}_2^B(\underline{\theta}) = \frac{\int_{\underline{\theta}}^a \Psi_2^B(x) dF_2(x)}{F_2(a)} = \frac{\underline{\theta} - a(1 - F_2(a))}{F_2(a)} < \underline{\theta}$. It follows that $\underline{w}(\underline{\theta})$ and $\bar{w}(\underline{\theta})$ exist, are unique, and decreasing respectively increasing, and that for $\underline{\theta} \geq 1$ and $r > 0$, inequality (A.1) holds if and only if $w \in [\underline{w}(\underline{\theta}), \bar{w}(\underline{\theta})]$. From the definitions, we have $\underline{w}(1) = \bar{w}(1) = 1/2$. ■

Proof of Theorem 2. For constant marginal values, given $\underline{\theta} \geq 0$ and $r \in [0, 1]$, let $\Pi_c(r)$ denote the expected budget surplus under ex post efficiency and binding IR for agents' worst-off types. It then follows that ex post efficiency is possible with $w = 1/2$ if and only if $\Pi_c(r) \geq 0$, where using (4), Π_c can be written as

$$\Pi_c(r) = \mathbb{E}_{\theta_1} [\Psi_1(\theta_1, \omega_1^e(r)) q_1^e(\theta_1)] + \mathbb{E}_{\theta_2} [\Psi_2(\theta_2, \omega_2^e(r)) q_2^e(\theta_2)] - r\omega_1^e(r) - (1 - r)\omega_2^e(r),$$

where $\omega_i^e(r)$ is agent i 's worst-off type under ex post efficiency. The following lemma will be useful:

Lemma A.1. *Given $\underline{\theta} \in [0, 1)$, $\Pi_c(r)$ is concave in r .*

Proof. Let $\omega_i^e(r)$ denote a worst-off type of agent i under the ex post efficient allocation rule and ownership $(r_1, r_2) = (r, 1 - r)$. To show that Π_c is concave, we show that for all $\lambda \in (0, 1)$

and $r, \hat{r} \in [0, 1]$ (dropping the subscript on Π_c to ease notation):

$$\Pi(\lambda r + (1 - \lambda)\hat{r}) \geq \lambda \Pi(r) + (1 - \lambda)\Pi(\hat{r}). \quad (\text{A.2})$$

Note that Π can be written as $\Pi(r) = \Pi_1(r) + \Pi_2(r)$, where

$$\Pi_i(r) \equiv \int_{\underline{\theta}_i}^{\omega_i^e(r)} \Psi_i^S(\theta) q_i^e(\theta) dF_i(\theta) + \int_{\omega_i^e(r)}^{\bar{\theta}_i} \Psi_i^B(\theta) q_i^e(\theta) dF_i(\theta) - r_i \omega_i^e(r).$$

Using the definition of $\Pi_i(r)$ and $(\Psi_i^S(\theta) - \Psi_i^B(\theta))f_i(\theta) = 1$, we have

$$\begin{aligned} \Pi_i(\lambda r + (1 - \lambda)\hat{r}) &= \lambda \Pi_i(r) + (1 - \lambda)\Pi_i(\hat{r}) \\ &\quad + \lambda \int_{\omega_i^e(r)}^{\omega_i^e(\lambda r + (1 - \lambda)\hat{r})} q_i^e(\theta) d\theta - (1 - \lambda) \int_{\omega_i^e(\lambda r + (1 - \lambda)\hat{r})}^{\omega_i^e(\hat{r})} q_i^e(\theta) d\theta \\ &\quad + \lambda r_i \omega_i^e(r) + (1 - \lambda)\hat{r}_i \omega_i^e(\hat{r}) - (\lambda r_i + (1 - \lambda)\hat{r}_i) \omega_i^e(\lambda r + (1 - \lambda)\hat{r}). \end{aligned} \quad (\text{A.3})$$

Let $\lambda \in (0, 1)$ and $r, \hat{r} \in [0, 1]$ with $0 \leq r < \hat{r} \leq 1$ be given. The possible worst-off types are for agent 1: (1) $\omega_1^e(r) < \omega_1^e(\hat{r})$ with $r = q_1^e(\omega_1^e(r)) < q_1^e(\omega_1^e(\hat{r})) = \hat{r}$; (2) $\omega_1^e(r) < \omega_1^e(\hat{r}) = 1$ with $r = q_1^e(\omega_1^e(r)) < q_1^e(\omega_1^e(\hat{r})) < \hat{r}$; and (3) $\omega_1^e(r) = \omega_1^e(\hat{r}) = 1$ with $q_1^e(1) < r < \hat{r}$. In cases 1 and 2, $\int_{\omega_1^e(r)}^{\omega_1^e(\lambda r + (1 - \lambda)\hat{r})} q_1^e(\theta) d\theta > r (\omega_1^e(\lambda r + (1 - \lambda)\hat{r}) - \omega_1^e(r))$ and $\int_{\omega_1^e(\lambda r + (1 - \lambda)\hat{r})}^{\omega_1^e(\hat{r})} q_1^e(\theta) d\theta \leq \hat{r} (\omega_1^e(\hat{r}) - \omega_1^e(\lambda r + (1 - \lambda)\hat{r}))$, which, using (A.3), imply that

$$\Pi_1(\lambda r + (1 - \lambda)\hat{r}) > \lambda \Pi_1(r) + (1 - \lambda)\Pi_1(\hat{r}),$$

and in case 3, this holds with equality.

For agent 2, we have: (1) $\omega_2^e(r) > \omega_2^e(\hat{r})$ with $1 - \hat{r} = q_2^e(\omega_2^e(\hat{r})) < q_2^e(\omega_2^e(r)) = 1 - r$; (2) $\underline{\theta} = \omega_2^e(\hat{r}) < \omega_2^e(r)$ with $1 - \hat{r} < q_2^e(\underline{\theta}) < q_2^e(\omega_2^e(r)) = 1 - r$; or (3) $\underline{\theta} = \omega_2^e(\hat{r}) = \omega_2^e(r)$ with $1 - \hat{r} < 1 - r < q_2^e(\underline{\theta})$. In cases 1 and 2, $\int_{\omega_2^e(\lambda r + (1 - \lambda)\hat{r})}^{\omega_2^e(r)} q_2^e(\theta) d\theta \leq (1 - r) (\omega_2^e(r) - \omega_2^e(\lambda r + (1 - \lambda)\hat{r}))$ and $\int_{\omega_2^e(\hat{r})}^{\omega_2^e(\lambda r + (1 - \lambda)\hat{r})} q_2^e(\theta) d\theta > (1 - \hat{r}) (\omega_2^e(\lambda r + (1 - \lambda)\hat{r}) - \omega_2^e(\hat{r}))$, which, using (A.3), imply that

$$\Pi_2(\lambda r + (1 - \lambda)\hat{r}) > \lambda \Pi_2(r) + (1 - \lambda)\Pi_2(\hat{r}),$$

and in case 3, this holds with equality.

Thus, $\Pi_i(\lambda r + (1 - \lambda)\hat{r}) \geq \lambda \Pi_i(r) + (1 - \lambda)\Pi_i(\hat{r})$ for each agent, implying that (A.2) holds and completing the proof. $\square$

The set of ownership structures consistent with ex post efficiency $\mathcal{R}_c(\underline{\theta})$ satisfies $\mathcal{R}_c(\underline{\theta}) = \{r \in [0, 1] \mid \Pi_c(r) \geq 0\}$. Extending existing results (Liu et al., 2026) to the case of agents with differing supports, we have the following, a version of which also holds for the multilateral bargaining setup:³⁷

³⁷One needs to replace $r^*(\underline{\theta}) \in [0, 1]$ by $\mathbf{r}^*(\underline{\theta}) \in [0, 1]^{n-1}$ with $\sum_{i=1}^{n-1} r_i^* \leq 1$ and adjust the definition of Π_c to apply for $n \geq 2$ agents.

Lemma A.2. *Given $\underline{\theta} \in [0, 1)$, there exists ownership $r^*(\underline{\theta}) \in [0, 1]$ that equalizes agents' worst-off types; moreover, $r^*(\underline{\theta})$ maximizes $\Pi_c(r)$ and $\Pi_c(r^*(\underline{\theta})) > 0$.*

Proof. See Online Appendix B.5, which proves the more general version for the multilateral bargaining setup.

Lemma A.2 implies that $r^*(\underline{\theta}) \in \mathcal{R}_c(\underline{\theta})$, which guarantees that $\mathcal{R}_c(\underline{\theta})$, and hence $\mathcal{E}_c(\underline{\theta})$, are nonempty for $\underline{\theta} \in [0, 1)$. They are also nonempty for $\underline{\theta} \geq 1$ because in that case ex post efficiency is possible for any $r \in [0, 1]$ and $w = 1/2$ (e.g., with a posted price mechanism with price $p \in [1, \underline{\theta}]$). This completes the proof that for any $\underline{\theta} \geq 0$, $\mathcal{E}_c(\underline{\theta})$ is nonempty.

We next prove that $r^*(\underline{\theta})$ is decreasing for $\underline{\theta} \in [0, 1)$.

Lemma A.3. *Ownership $r^*(\underline{\theta})$ defined in Lemma A.2 is decreasing in $\underline{\theta}$ for $\underline{\theta} \in [0, 1)$.*

Proof. Assume that $\underline{\theta} \in [0, 1)$. From Lemma A.2, r^* equalizes agents' worst-off types, implying that there exists ω^* such that $q_1^e(\omega^*) = r^*$ and $q_2^e(\omega^*) = 1 - r^*$. Explicitly noting the reliance of q_1^e on $\underline{\theta}$, ω^* is defined by $q_1^e(\omega^*, \underline{\theta}) = 1 - q_2^e(\omega^*)$, where $q_1^e(\theta_1, \underline{\theta}) = \int_0^1 Q_1(\theta_1, x + \underline{\theta}) dF_2^P(x)$ and $q_2^e(\theta_2) = \int_0^1 Q_2(\theta_1, \theta_2) dF_1(\theta_1)$. Because an agent's worst-off type must be in its type support, $\omega^* \in [\underline{\theta}, 1]$, which implies that $q_1^e(\omega^*, \underline{\theta}) = F_2^P(\omega^* - \underline{\theta})$ and $q_2^e(\omega^*) = F_1(\omega^*)$. Using these expressions and $q_1^e(\omega^*, \underline{\theta}) + q_2^e(\omega^*) = 1$ and totally differentiating with respect to $\underline{\theta}$, we have $\frac{\partial \omega^*}{\partial \underline{\theta}} = \frac{f_2^P(\omega^* - \underline{\theta})}{f_2^P(\omega^* - \underline{\theta}) + f_1(\omega^*)} > 0$. Because $q_2^e(\omega^*) = 1 - r^*$ and $q_2^e(\theta)$ is increasing for $\theta \in [\underline{\theta}, 1]$, it follows that $\frac{\partial r^*}{\partial \underline{\theta}} < 0$. $\square$

Combining Lemmas A.1, A.2, and A.3, it follows that for $\underline{\theta} \in [0, 1)$, there exist $\underline{r}(\underline{\theta})$ and $\bar{r}(\underline{\theta})$ such that $0 \leq \underline{r}(\underline{\theta}) < r^*(\underline{\theta}) < \bar{r}(\underline{\theta}) \leq 1$ and such that for all $r \in [\underline{r}(\underline{\theta}), \bar{r}(\underline{\theta})]$, $\Pi_c(r) \geq 0$. For $\underline{\theta} \in [0, 1)$, Myerson and Satterthwaite (1983) show that ex post efficiency is impossible with extremal ownership, implying that neither 0 nor 1 are elements of $\mathcal{R}_c(\underline{\theta})$, so we have $0 < \underline{r}(\underline{\theta}) < r^*(\underline{\theta}) < \bar{r}(\underline{\theta}) < 1$.

Because an agent's worst-off type must be in its type support, in the limit as $\underline{\theta}$ goes to 1 from below, equalized worst-off types approach 1. Using a result first established by Cramton et al. (1987), agent 2's worst-off type $\hat{\theta}_2$ must satisfy either (a) $q_2^e(\hat{\theta}_2) = 1 - r$ or (b) $\hat{\theta}_2 = \underline{\theta}$ and $q_2(\underline{\theta}) > 1 - r$. Thus, noting that $q_2^e(1) = F_1(1) = 1$, a worst-off type of 1 requires that $r = 0$. Further, when r approaches zero and $\underline{\theta}$ approaches 1, the expected net payoff, and hence maximized budget surplus, approaches zero, with the implication that the set of ownerships inducing positive revenue under ex post efficiency approaches the singleton set $\{0\}$.

Turning to the results on bargaining weights, recall that $\tau_c^* \equiv \max\{\Psi_1^S(1), 1 - \Psi_2^{B,P}(0)\}$. **Range (1):** For $\underline{\theta} \in [0, 1]$, ex post efficiency is possible for $w = 1/2$ and $r \in \mathcal{R}_c(\underline{\theta})$, as just seen, so we are left to show that for $w \neq 1/2$, it is not possible. To do this, it suffices to recall

the allocation rule of Proposition 1. This rule is not ex post efficient if $w \neq 1/2$ because then the ironed weighted virtual type functions of the two agents differ (see also Lemma B.10).

Range (2): For $\underline{\theta} \in (1, \tau_c^*)$, ex post efficiency is possible for $w = 1/2$ for any r . In addition, with $\underline{\theta} > 1$, small enough departures from equal bargaining weights only affect the ironed weighted virtual type functions in such a way that agent 2's ironed weighted virtual type is still larger than agent 1's for all possible θ_2 and θ_1 . The larger is $\underline{\theta}$, the larger can these departures from equality be, which proves that $\mathcal{W}_c(\underline{\theta})$ increases in the set inclusion sense. In addition, the monotonicity of the virtual type functions in w_i imply that $\mathcal{W}_c(\underline{\theta})$ is convex, giving us that $\underline{w}(\underline{\theta})$ is decreasing and $\overline{w}(\underline{\theta})$ is increasing, which completes the proof of (iii). In addition, $0 \leq \underline{w}(\underline{\theta}) < \overline{w}(\underline{\theta}) \leq 1$ for $\underline{\theta} > 1$. For $\underline{\theta} < \tau_c^*$, either $\Psi_1^S(1) > \underline{\theta}$, implying that trade will not be ex post efficient for $w = 0$ or sufficiently close to 0, or $\underline{\theta} < 1 - \Psi_1^B(0)$, implying that trade will not be ex post efficient for $w = 1$ or sufficiently close to 1 (or both). Thus, for $\underline{\theta} \in [1, \tau_c^*)$, $0 \leq \underline{w}(\underline{\theta}) < \overline{w}(\underline{\theta}) \leq 1$, with at least one weak inequality strict, which proves the first part of (iv). For additional characterization of $\underline{w}$ and $\overline{w}$, see Proposition 7 and its proof.

Range (3): For $\underline{\theta} \geq \tau_c^*$, agent 2 is always a buyer and agent 1 always a seller of r units. Because the weighted virtual type functions are monotone in w and ex post efficiency obtains for any $w \in [0, 1]$, we have $\underline{w}(\underline{\theta}) = 0$ and $\overline{w}(\underline{\theta}) = 1$. By continuity, we have the limit result.

Combining these results, we obtain the characterization of $\mathcal{E}_c(\underline{\theta})$, completing the proof of Theorem 2. ■

Proof of Proposition 3. Here we prove part (i) of the proposition, which then implies part (ii). See Online Appendix B.6 for the proof of part (iii).

Concavity follows from the slope of the frontier, which we now prove. Let $\langle \mathbf{Q}^w, \mathbf{M}^w \rangle$ be the bargaining mechanism for a given w . Letting $u_i^w(\theta_i)$ denote agent i 's interim expected net payoff given w , away from ex post efficiency, the expected net payoff frontier is given by $(\mathbb{E}_{\theta_1}[u_1^w(\theta_1)], \mathbb{E}_{\theta_2}[u_2^w(\theta_2)])_{w \in [0,1]}$, where $\frac{d\mathbb{E}_{\theta_1}[u_1^w(\theta_1)]}{dw} > 0$, and so the frontier has slope $\frac{d\mathbb{E}_{\theta_2}[u_2^w(\theta_2)]}{dw} / \frac{d\mathbb{E}_{\theta_1}[u_1^w(\theta_1)]}{dw}$. By the envelope theorem, the derivative with respect to w of the optimized objective for the bargaining problem satisfies

$$\begin{aligned} & \frac{d}{dw} \left(\mathbb{E}_{\theta} [w(Q_1^w(\theta)\theta_1 - M_1^w(\theta)) + (1-w)(Q_2^w(\theta)\theta_2 - M_2^w(\theta))] \right) \\ &= \mathbb{E}_{\theta} [Q_1^w(\theta)\theta_1 - M_1^w(\theta) - (Q_2^w(\theta)\theta_2 - M_2^w(\theta))]. \end{aligned} \quad (\text{A.4})$$

Using $u_i^w(\theta_i) = (q_i^w(\theta_i) - r_i)\theta_i - m_i^w(\theta_i)$, we can write $w\mathbb{E}_{\theta_1}[u_1^w(\theta_1)] + (1-w)\mathbb{E}_{\theta_2}[u_2^w(\theta_2)] = \mathbb{E}_{\theta} [w(Q_1^w(\theta)\theta_1 - M_1^w(\theta)) + (1-w)(Q_2^w(\theta)\theta_2 - M_2^w(\theta))] - \mathbb{E}_{\theta} [wr\theta_1 + (1-w)(1-r)\theta_2]$. Dif-

ferentiating this equation with respect to w and using (A.4), we get

$$\begin{aligned}
& \mathbb{E}_{\theta_1}[u_1^w(\theta_1)] - \mathbb{E}_{\theta_2}[u_2^w(\theta_2)] + w \frac{d\mathbb{E}_{\theta_1}[u_1^w(\theta_1)]}{dw} + (1-w) \frac{d\mathbb{E}_{\theta_2}[u_2^w(\theta_2)]}{dw} \\
= & \mathbb{E}_{\theta}[Q_1^w(\theta)\theta_1 - M_1(\theta) - (Q_2^w(\theta)\theta_2 - M_2(\theta))] - \mathbb{E}_{\theta}[r\theta_1 - (1-r)\theta_2] \\
= & \mathbb{E}_{\theta_1}[u_1^w(\theta_1)] - \mathbb{E}_{\theta_2}[u_2^w(\theta_2)],
\end{aligned}$$

which gives $\frac{d\mathbb{E}_{\theta_2}[u_2^w(\theta_2)]}{dw} / \frac{d\mathbb{E}_{\theta_1}[u_1^w(\theta_1)]}{dw} = -\frac{w}{1-w}$. In contrast, when ex post efficiency is achieved, $\mathbb{E}_{\theta_1}[u_1^w(\theta_1)] + \mathbb{E}_{\theta_2}[u_2^w(\theta_2)]$ is constant, which implies that the slope of the frontier at ex post efficiency is -1 . This completes the proof of part (i) of the proposition. ■

References

- AUSTRALIAN COMPETITION & CONSUMER COMMISSION (2018): “ACCC Dairy Inquiry: Final Report,” <https://www.accc.gov.au/about-us/publications/dairy-inquiry-final-report>.
- AUSUBEL, L. M., P. CRAMTON, AND R. J. DENECKERE (2002): “Bargaining with Incomplete Information,” in *Handbook of Game Theory*, ed. by R. Aumann and S. Hart, Elsevier Science B.V., vol. 3, 1897–1945.
- AVIGNON, R., C. CHAMBOLLE, E. GUIGUE, AND H. MOLINA (2025): “Markups, Markdowns, and Bargaining in a Vertical Supply Chain,” Working Paper.
- BACKUS, M., T. BLAKE, B. LARSEN, AND S. TADELIS (2020): “Sequential Bargaining in the Field: Evidence from Millions of Online Bargaining Interactions,” *Quarterly Journal of Economics*, 135, 1319–1361.
- BACKUS, M., T. BLAKE, J. PETTUS, AND S. TADELIS (forth.): “Communication, Learning, and Bargaining Breakdown: An Empirical Analysis,” *Management Science*.
- BACKUS, M., T. BLAKE, AND S. TADELIS (2019): “On the Empirical Content of Cheap-Talk Signaling: An Application to Bargaining,” *Journal of Political Economy*, 127, 1599–1628.
- BARKLEY, A., D. GENESOVE, AND J. HANSEN (2025a): “Auction Identification with Unobserved Rejected Offers,” Working Paper, Hebrew University of Jerusalem.
- (2025b): “Haggle or Hammer? Dual-Mechanism Housing Search,” Working Paper, Hebrew University of Jerusalem.
- BLEAKLEY, H. AND J. FERRIE (2014): “Land Openings on the Georgia Frontier and the Coase Theorem in the Short and Long Run,” University of Chicago.
- BOYD, S. AND L. VANDENBERGHE (2004): *Convex Optimization*, Cambridge University Press.
- BULOW, J. AND J. ROBERTS (1989): “The Simple Economics of Optimal Auctions,” *Journal of Political Economics*, 97, 1060–1090.
- BYRNE, D. P., L. A. MARTIN, AND J. S. NAH (2022): “Price Discrimination, Search, and Negotiation in an Oligopoly: A Field Experiment in Retail Electricity,” *Quarterly Journal of Economics*, 137, 2499–2537.
- CHATTERJEE, K. AND W. SAMUELSON (1983): “Bargaining under Incomplete Information,” *Operations Research*, 31, 835–851.
- CHE, Y.-K. (2006): “Beyond the Coasian Irrelevance: Asymmetric Information,” Unpublished Lecture Notes, Columbia University.
- CHONÉ, P., L. LINNEMER, AND T. VERGÉ (2024): “Double Marginalization, Market Foreclosure, and Vertical Integration,” *Journal of the European Economic Association*, 22, 1884–1935.
- COASE, R. H. (1960): “The Problem of Social Cost,” *Journal of Law & Economics*, 3, 1–44.
- CRAMTON, P., R. GIBBONS, AND P. KLEMPERER (1987): “Dissolving a Partnership Efficiently,” *Econometrica*, 55, 615–632.
- DEMIRER, M. AND M. RUBENS (2025): “Welfare Effects of Buyer and Seller Power,” NBER Working Paper # 33371.
- FIGUEROA, N. AND V. SKRETA (2012): “Asymmetric Partnerships,” *Econ. Letters*, 115, 268–271.
- FOWLIE, M. AND J. M. PERLOFF (2013): “Distributing Pollution Rights in Cap-and-Trade Programs: Are Outcomes Independent of Allocation?” *Review of Economics and Stat.*, 95, 1640–1652.

- FRIEDEN, R., K. JAYAKAR, AND E.-A. PARK (2020): “There’s Probably a Blackout in Your Television Future: Tracking New Carriage Negotiation Strategies Between Video Content Programmers and Distributors,” *Columbia Journal of Law and the Arts*, 43, 487–515.
- GALBRAITH, J. K. (1952): *American Capitalism: The Concept of Countervailing Power*, Houghton Mifflin, Boston.
- (1954): “Countervailing Power,” *American Economic Review (P&P)*, 44, 1–6.
- GRESIK, T. AND M. SATTERTHWAIT (1989): “The Rate at which a Simple Market Converges to Efficiency as the Number of Traders Increases: An Asymptotic Result for Optimal Trading Mechanisms,” *Journal of Economic Theory*, 48, 304–332.
- GROSSMAN, S. J. AND O. D. HART (1986): “The Costs and Benefits of Ownership: A Theory of Vertical and Lateral Integration,” *Journal of Political Economy*, 94, 691–719.
- HART, O. AND J. MOORE (1990): “Property Rights and the Nature of the Firm,” *Journal of Political Economy*, 98, 1119–1158.
- HO, K. AND R. LEE (2017): “Insurer Competition in Health Care Markets,” *Econometrica*, 85, 379–417.
- KRÄHMER, D. AND R. STRAUZ (2007): “VCG Mechanisms and Efficient Ex Ante Investments with Externalities,” *Economics Letters*, 94, 192–196.
- KRUEGER, A. B. (2018): “Luncheon Address: Reflections on Dwindling Worker Bargaining Power and Monetary Policy,” August 24, 2018, Luncheon Address at the Jackson Hole Econ. Symposium.
- LARSEN, B. J. (2021): “The Efficiency of Real-World Bargaining: Evidence from Wholesale Used-Auto Auctions,” *Review of Economic Studies*, 88, 851–882.
- LARSEN, B. J., C. H. LU, AND A. L. ZHANG (2024): “Intermediaries in Bargaining: Evidence from Business-to-Business Used-Car Inventory Negotiations,” Working Paper, Wash. U. St. Louis.
- LARSEN, B. J. AND A. ZHANG (2025): “Quantifying Bargaining Power Under Incomplete Information: A Supply-Side Analysis of the Used-Car Industry,” Working Paper, Wash. U. St. Louis.
- LE, T. (2024): “(In)complete Information Bargaining,” Working Paper, University of Melbourne.
- LEE, R. S., M. D. WHINSTON, AND A. YURUKOGLU (2021): “Structural Empirical Analysis of Contracting in Vertical Markets,” in *Handbook of Industrial Org.*, Elsevier, vol. 4, 673–742.
- LIU, B., S. LOERTSCHER, AND L. M. MARX (2026): “Efficient Consignment Auctions,” *Review of Economics and Statistics*, 108, 225–240.
- LOERTSCHER, S. AND L. M. MARX (2019): “Merger Review for Markets with Buyer Power,” *Journal of Political Economy*, 127, 2967–3017.
- (2022): “Incomplete Information Bargaining with Applications to Mergers, Investment, and Vertical Integration,” *American Economic Review*, 112, 616–649.
- (2023): “Asymptotically Optimal Prior-Free Asset Market Mechanisms,” *Games and Economic Behavior*, 137, 68–90.
- (2024): “Mergers, Remedies, and Incomplete Information,” Working Paper, U. Melbourne.
- LOERTSCHER, S. AND E. V. MUIR (2025): “Optimal Hotelling Auctions,” Working Paper.
- LOERTSCHER, S. AND A. NIEDERMAYER (2019): “An Optimal Pricing Theory of Transaction Fees in Thin Markets,” Working Paper, University of Melbourne.
- LOERTSCHER, S. AND C. WASSER (2019): “Optimal Structure and Dissolution of Partnerships,”

- Theoretical Economics*, 14, 1063–1114.
- LU, H. AND J. ROBERT (2001): “Optimal Trading Mechanisms with Ex Ante Unidentified Traders,” *Journal of Economic Theory*, 97, 50–80.
- MAKOWSKI, L. AND C. MEZZETTI (1993): “The Possibility of Efficient Mechanisms for Trading and Indivisible Object,” *Journal of Economic Theory*, 59, 451–465.
- MILGROM, P. (2004): *Putting Auction Theory to Work*, Cambridge University Press.
- MILGROM, P. AND I. SEGAL (2002): “Envelope Theorems for Arbitrary Choice Sets,” *Econometrica*, 70, 583–601.
- MUSSA, M. AND S. ROSEN (1978): “Monopoly and Product Quality,” *Journal of Economic Theory*, 18, 301–317.
- MYERSON, R. B. (1981): “Optimal Auction Design,” *Mathematics of Operations Research*, 6, 58–73.
- MYERSON, R. B. AND M. A. SATTERTHWAITE (1983): “Efficient Mechanisms for Bilateral Trading,” *Journal of Economic Theory*, 29, 265–281.
- MYLOVANOV, T. AND T. TRÖGER (2014): “Mechanism Design by an Informed Principal: Private Values with Transferable Utility,” *Review of Economic Studies*, 81, 1668–1707.
- NASH, J. F. (1950): “The Bargaining Problem,” *Econometrica*, 18, 155–162.
- OECD (2007): “Competition and Efficient Usage of Payment Cards 2006,” *OECD Policy Roundtables*, DAF/COMP(2006)32.
- (2011): “Competition and Efficient Usage of Payment Cards,” *OECD Journal: Competition Law & Policy*, 2009.
- SHAPLEY, L. S. (1951): “Notes on the n -Person Game – II: The Value of an n -Person Game,” RAND Corporation RM 670.
- STEPTOE, M. L. (1993): “The Power-Buyer Defense in Merger Cases,” *Antitrust L. J.*, 61, 493–504.
- STIGLER, G. J. (1954): “The Economist Plays with Blocs,” *American Economic Review (Papers and Proceedings)*, 44, 7–14.
- U.S. DEPARTMENT OF JUSTICE AND THE FEDERAL TRADE COMMISSION (1996): “Statement of Antitrust Enforcement Policy in Health Care,” <https://www.justice.gov/atr/public/guidelines/1791.htm>.
- UYANIK, M. AND D. YENGİN (2023): “Expropriation Power in Private Dealings: Quota Rule in Collective Sales,” *Games and Economic Behavior*, 141, 548–580.
- VAN DER MAELEN, S., E. BREUGELMANS, AND K. CLEEREN (2017): “The Clash of the Titans: On Retailer and Manufacturer Vulnerability in Conflict Delistings,” *J. of Marketing*, 81, 118–135.
- VICKREY, W. (1961): “Counterspeculation, Auction, and Competitive Sealed Tenders,” *Journal of Finance*, 16, 8–37.
- WILLIAMS, S. R. (1987): “Efficient Performance in Two Agent Bargaining,” *Journal of Economic Theory*, 41, 154–172.
- WILSON, R. (1993): *Nonlinear Pricing*, Oxford University Press.
- XIAO, J. (2018): “Bargaining Orders in a Multi-Person Bargaining Game,” *Games and Economic Behavior*, 364–379.

Online Appendix

to accompany

“Complete and incomplete information bargaining in a model
with independent private values”

by

Simon Loertscher and Leslie M. Marx

January 30, 2026

Sections in this Online Appendix are labeled starting with “B” to distinguish these appendices from Appendix A in the paper.

B Omitted proofs

B.1 Equivalence of BIC and DIC (footnote 17)

Lemma B.1. *With both constant and decreasing marginal values, Bayesian incentive compatibility (BIC) and dominant-strategy incentive compatibility (DIC) are equivalent, that is, given BIC and interim IR mechanism $\langle \mathbf{Q}, \mathbf{M}' \rangle$, there exists $\mathbf{M}$ such that $\langle \mathbf{Q}, \mathbf{M} \rangle$ satisfies DIC and interim IR and has the same expected budget surplus under binding IR for agents’ worst-off types as does $\langle \mathbf{Q}, \mathbf{M}' \rangle$.*

Proof. Suppose that $\langle \mathbf{Q}, \mathbf{M}' \rangle$ satisfies BIC and interim IR. We show that there exists $\mathbf{M}$ such that $\langle \mathbf{Q}, \mathbf{M} \rangle$ satisfies DIC and interim IR, and such that both mechanisms have the same expected budget surplus under binding IR for agents’ worst-off types.

Define $q_i(\theta_i) \equiv \mathbb{E}_{\boldsymbol{\theta}_{-i}}[Q_i(\theta_i, \boldsymbol{\theta}_{-i})]$. Letting ω_i denote agent i ’s interim worst-off type, we have (see Cramton et al., 1987):

$$\omega_i = \begin{cases} \underline{\theta}_i & \text{if } r_i < q_i(\underline{\theta}_i), \\ \omega_i \text{ s.t. } q_i(\omega_i) = r_i & \text{if } r_i \in [q_i(\underline{\theta}_i), q_i(\bar{\theta}_i)] \text{ and } \exists \omega_i \text{ s.t. } q_i(\omega_i) = r_i, \\ \omega_i \text{ s.t. } \lim_{\theta \uparrow \omega_i} q_i(\theta) < r_i \text{ and } \lim_{\theta \downarrow \omega_i} q_i(\theta) > r_i & \text{if } r_i \in [q_i(\underline{\theta}_i), q_i(\bar{\theta}_i)] \text{ and } \nexists \omega_i \text{ s.t. } q_i(\omega_i) = r_i, \\ \bar{\theta}_i & \text{if } r_i > q_i(\bar{\theta}_i). \end{cases}$$

Recall that by our characterization of the BIC, the interim expected payment rule under BIC is $\mathbb{E}_{\boldsymbol{\theta}_{-i}}[M'_i(\theta_i, \boldsymbol{\theta}_{-i})] \equiv \mathbb{E}_{\boldsymbol{\theta}_{-i}}[V(\theta_i, Q_i(\boldsymbol{\theta}))] - V(\theta_i, r_i) - \int_{\omega_i}^{\theta_i} (q_i(x, \boldsymbol{\theta}_{-i}) - r_i) dx - u_i(\omega_i)$. We show that the mechanism $\langle \mathbf{Q}, \mathbf{M} \rangle$, where

$$M_i(\theta_i, \boldsymbol{\theta}_{-i}) \equiv V(\theta_i, Q_i(\boldsymbol{\theta})) - V(\theta_i, r_i) - \int_{\omega_i}^{\theta_i} (Q_i(x, \boldsymbol{\theta}_{-i}) - r_i) dx - u_i(\omega_i),$$

satisfies DIC and interim IR. To do so, define agent 1's payoff net of its outside option if types are $(\theta_i, \boldsymbol{\theta}_{-i})$ and firm 1 reports θ'_i as

$$U_i(\theta_i, \boldsymbol{\theta}_{-i}; \theta'_i) \equiv V(\theta_i, Q_i(\theta'_i, \boldsymbol{\theta}_{-i})) - V(\theta_i, r_i) - M_i(\theta'_i, \boldsymbol{\theta}_{-i}).$$

First, we show that DIC holds, i.e., that for all θ_i , θ'_i , and $\boldsymbol{\theta}_{-i}$, we have $U_i(\theta_i, \boldsymbol{\theta}_{-i}; \theta_i) \geq U_i(\theta_i, \boldsymbol{\theta}_{-i}; \theta'_i)$. To see this, note that

$$\begin{aligned} U_i(\theta_i, \boldsymbol{\theta}_{-i}; \theta_i) - U_i(\theta_i, \boldsymbol{\theta}_{-i}; \theta'_i) &= V(\theta_i, Q_i(\theta_i, \boldsymbol{\theta}_{-i})) - V(\theta_i, r_i) - M_i(\theta_i, \boldsymbol{\theta}_{-i}) \\ &\quad - V(\theta_i, Q_i(\theta'_i, \boldsymbol{\theta}_{-i})) + V(\theta_i, r_i) + M_i(\theta'_i, \boldsymbol{\theta}_{-i}) \\ &= \int_{\theta'_i}^{\theta_i} Q_i(x, \boldsymbol{\theta}_{-i}) dx - (\theta_i - \theta'_i) Q_i(\theta'_i, \boldsymbol{\theta}_{-i}) \\ &\geq 0, \end{aligned}$$

where the inequality uses that $Q_i(\theta_i, \boldsymbol{\theta}_{-i})$ is nondecreasing in θ_i .

Second, we show that interim IR holds, i.e., $\mathbb{E}_{\boldsymbol{\theta}_{-i}}[U_i(\theta_i, \boldsymbol{\theta}_{-i}; \theta_i)] \geq 0$. To see this, note that

$$\begin{aligned} \mathbb{E}_{\boldsymbol{\theta}_{-i}}[U_i(\theta_i, \boldsymbol{\theta}_{-i}; \theta_i)] &= \mathbb{E}_{\boldsymbol{\theta}_{-i}}[V(\theta_i, Q_i(\theta_i, \boldsymbol{\theta}_{-i})) - V(\theta_i, r_i) - M_i(\theta_i, \boldsymbol{\theta}_{-i})] \\ &= \int_{\omega_i}^{\theta_i} (q_i(x) - r_i) dx + u_i(\omega_i) \\ &= u_i(\omega_i) + \begin{cases} \int_{\underline{\theta}_i}^{\theta_i} (q_i(x) - r_i) dx & \text{if } r_i < q_i(\underline{\theta}_i) \\ \int_{\omega_i}^{\theta_i} (q_i(x) - q_i(\omega_i)) dx & \text{if } r_i \in [q_i(\underline{\theta}_i), q_i(\bar{\theta}_i)] \text{ and } q_i(\omega_i) = r_i, \\ \int_{\omega_i}^{\theta_i} (q_i(x) - r_i) dx & \text{if } r_i \in [q_i(\underline{\theta}_i), q_i(\bar{\theta}_i)] \text{ and } q_i(\omega_i) \neq r_i, \\ \int_{\theta_i}^{\bar{\theta}_i} (r_i - q_i(x)) dx & \text{if } r_i > q_i(\bar{\theta}_i) \end{cases} \\ &\geq u_i(\omega_i), \end{aligned}$$

where the inequality uses that $q_i(\theta_i)$ is nondecreasing in θ_i and, for the third row in the bracketed expression, that $r_i < q_i(\theta_i)$ for all $\theta_i > \omega_i$ and $r_i > q_i(\theta_i)$ for all $\theta_i < \omega_i$.

Thus, $\langle \mathbf{Q}, \mathbf{M} \rangle$ satisfies DIC and interim IR. Further, $\mathbb{E}_{\boldsymbol{\theta}_{-i}}[M_i(\theta_i, \boldsymbol{\theta}_{-i})] = \mathbb{E}_{\boldsymbol{\theta}_{-i}}[M'_i(\theta_i, \boldsymbol{\theta}_{-i})]$, and so the expected budget surplus under $\langle \mathbf{Q}, \mathbf{M}' \rangle$ is the same as under $\langle \mathbf{Q}, \mathbf{M} \rangle$. ■

B.2 No benefit from using randomization ex post (footnote 25)

With decreasing marginal values, an agent's utility function is concave in the quantity and, therefore, exhibits a form of risk aversion. In principle, the designer could leverage this to its benefit by using randomization ex post to make the agents' incentive compatibility constraints slacker. However, the following lemma shows that this is not the case. (With constant marginal values, the agents are risk neutral and thus the issue is moot in that case.)

Lemma B.2. *With decreasing marginal values, there is no benefit from using randomization ex post in the bargaining mechanism.*

Proof. Standard arguments imply that IC (i.e., BIC) requires q_i to be monotone and that the payment rule is such that the mechanism satisfies IC. So for IC, volatility plays no role—all that matters are the interim expected allocations q_i . Moreover, from a revenue perspective nothing good can come from ex post randomization because increasing variance decreases revenue. To be more precise, denoting by

$$\sigma_{Q_i}(\theta) \equiv \mathbb{E}_{\theta_j}[(Q_i(\theta, \theta_j) - q_i(\theta))^2] = \mathbb{E}_{\theta_j}[Q_i(\theta, \theta_j)^2] - q_i(\theta)^2$$

the expected variance of i 's allocation given its type θ , the expected payment of i when of type θ can be written as

$$m_i(\theta) = \theta q_i(\theta) - \frac{\sigma_{Q_i}(\theta) + q_i(\theta)^2}{2} - \int_{\hat{\theta}}^{\theta} q_i(y) dy - (\theta r_i - r_i^2/2) - u_i(\hat{\theta}),$$

which is decreasing in $\sigma_{Q_i}(\theta)$. ■

B.3 Continuity and monotonicity of expected revenue in ρ (footnote 27)

B.3.1 Outline and overview

For a given parameterization of the problem $(\mathbf{w}, \mathbf{r})$, where w_i and r_i are i 's bargaining weight and ownership share, a given ρ , and a given type profile $\boldsymbol{\theta}$, agent i 's ex post allocation under the optimal mechanism, $Q_i(\boldsymbol{\theta}; \rho)$, will be either 0, 1, or some $\beta_i \in (0, 1)$.¹ Even though neither the ex post allocations $Q_i(\boldsymbol{\theta}; \rho)$ nor the virtual type functions

$$\Psi_i(\theta_i, \omega_i^*(\rho)) = \begin{cases} \theta_i + \frac{F_i(\theta_i)}{f_i(\theta_i)} & \theta_i \leq \omega_i^*(\rho), \\ \theta_i - \frac{1-F_i(\theta_i)}{f_i(\theta_i)} & \theta_i > \omega_i^*(\rho), \end{cases}$$

are continuous in ρ , one can show that expected revenue is continuous in ρ because:

- (i) the worst-off types $\omega_i^*(\rho)$ are continuous;
- (ii) $Q_i(\boldsymbol{\theta}; \rho)$ and $\Psi_i(\theta_i, \omega_i^*(\rho))$ are continuous (indeed mostly constant) in ρ on the interior of the subsets of types over which one integrates; and
- (iii) the boundaries of the integrals vary continuously in ρ .

¹As is well known, with two or more agents, one could dispense with ex post randomization (Chen et al., 2019). However, this does not change the fact that allocation rules with ex post randomization are optimal, and working with ex post randomization makes it easier to establish the desired continuity properties.

The boundaries where the allocations $Q_i(\boldsymbol{\theta}; \rho)$ change discontinuously are determined by comparisons of the ironed weighted virtual types $\bar{\Psi}_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i^*(\rho))$ and $\bar{\Psi}_{j, \frac{w_j}{\rho}}(\theta_j, \omega_j^*(\rho))$, where $j \neq i$ is the agent with the highest ironed weighted virtual type other than i . Continuity of these boundaries follows partly from the continuity of $\omega_i^*(\rho)$, i.e., point (i) above. Likewise, the expected revenue from agent i ,

$$\Pi_i(\rho) \equiv \mathbb{E}_{\boldsymbol{\theta}} [\Psi_i(\theta_i, \omega_i^*(\rho)) Q_i(\boldsymbol{\theta}; \rho)] - r_i \omega_i^*(\rho),$$

can be rewritten as

$$\Pi_i(\rho) \equiv \int_{\underline{\theta}_i}^{\omega_i^*(\rho)} \left(\theta_i + \frac{F_i(\theta_i)}{f_i(\theta_i)} \right) q_i(\theta_i; \rho) dF_i(\theta_i) + \int_{\omega_i^*(\rho)}^{\bar{\theta}_i} \left(\theta_i - \frac{1 - F_i(\theta_i)}{f_i(\theta_i)} \right) q_i(\theta_i; \rho) dF_i(\theta_i) - r_i \omega_i^*(\rho),$$

where $q_i(\theta_i; \rho) = \mathbb{E}_{\boldsymbol{\theta}_{-i}} [Q_i(\boldsymbol{\theta}; \rho)]$ is the interim expected allocation of agent i . The continuity of the bounds of integration and the continuity of $r_i \omega_i^*(\rho)$ are then, again, implications of (i). The key to establishing continuity of $\omega_i^*(\rho)$ is showing that the ironing parameters $z_i^*(\rho)$ (the level at which agent i 's ironed weighted virtual type, $\bar{\Psi}_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i^*(\rho))$, becomes flat) are continuous in ρ .

B.3.2 Formal statement and proof

Proposition B.1. *Expected revenue under binding IR for agents' worst-off types,*

$$\mathbb{E}_{\boldsymbol{\theta}} \left[\sum_{i=1}^n \Psi_i(\theta_i, \omega_i^*(\rho)) Q_i(\boldsymbol{\theta}; \rho) \right] - \sum_{i=1}^n r_i \omega_i^*(\rho),$$

is continuous in ρ .

Remark: We prove the proposition for the case of $n = 2$ and monotone virtual cost and virtual value functions. At the cost of additional notation and case distinctions, the arguments extend to more than two agents and nonmonotone virtual costs and virtual values. We discuss after the proof which arguments and features need to be adjusted and in which ways. There we also briefly sketch how the arguments extend to the case of decreasing marginal values.

Proof. We establish continuity of $\Pi_i(\rho) \equiv \mathbb{E}_{\boldsymbol{\theta}} [\Psi_i(\theta_i, \omega_i^*(\rho)) Q_i(\boldsymbol{\theta}; \rho)] - r_i \omega_i^*(\rho)$, which then implies continuity of the sum. Without loss of generality, we assume that $\max \mathbf{w} = 1$ so that $\rho \in [1, \infty)$. To begin, we establish the continuity of $z_i^*(\rho)$, $\omega_i^*(\rho)$, and $\beta_i(\rho)$ (if tie-breaking is used).

We begin with the following lemma:

Lemma B.3. For $n = 2$, monotone virtual cost and virtual value functions, and overlapping supports (i.e., $\underline{\theta}_1 \leq \underline{\theta}_2 < \bar{\theta}_1 \leq \bar{\theta}_2$) (see Lemma B.4 for the case of nonoverlapping supports):

$$(z_1^*(\rho), z_2^*(\rho)) = \begin{cases} \left(\Psi_{2, \frac{w_2}{\rho}}^S(F_2^{-1}(r)), \Psi_{1, \frac{w_1}{\rho}}^B(F_1^{-1}(1-r)) \right) & \text{if } r \in [0, \hat{r}(\rho)], \\ (z^*(\rho), z^*(\rho)) & \text{if } r \in (\hat{r}(\rho), \tilde{r}(\rho)), \\ \left(\Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(r)), \Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1-r)) \right) & \text{if } r \in [\tilde{r}(\rho), 1], \end{cases}$$

where $z^*(\rho)$ is defined by (B.2) and $\hat{r}(\rho)$ and $\tilde{r}(\rho)$ with $0 < \hat{r}(\rho) \leq \tilde{r}(\rho) < 1$ are defined by

$$\Psi_{2, \frac{w_2}{\rho}}^S(F_2^{-1}(\hat{r})) = \Psi_{1, \frac{w_1}{\rho}}^B(F_1^{-1}(1 - \hat{r})) \quad \text{and} \quad \Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - \tilde{r})) = \Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(\tilde{r})),$$

where $\hat{r}'(\rho) < 0$ and $\tilde{r}'(\rho) > 0$. Further,

$$\beta^*(\rho) = \frac{r - F_2\left(\Psi_{2, \frac{w_2}{\rho}}^{S^{-1}}(z^*(\rho))\right)}{F_2\left(\Psi_{2, \frac{w_2}{\rho}}^{B^{-1}}(z^*(\rho))\right) - F_2\left(\Psi_{2, \frac{w_2}{\rho}}^{S^{-1}}(z^*(\rho))\right)}. \quad (\text{B.1})$$

Before proving Lemma B.3, note that because the expressions and bounds in the expression for $(z_1^*(\rho), z_2^*(\rho))$ in Lemma B.3 all vary continuously with ρ , it follows that $(z_1^*(\rho), z_2^*(\rho))$ is continuous (indeed, differentiable) in ρ . Thus, Lemma B.3 has the following corollary:

Corollary 2. For $n = 2$, monotone virtual cost and virtual value functions, and overlapping supports, $z_i^*(\rho)$, and hence $\omega_i^*(\rho)$ and $\beta^*(\rho)$, are continuous in ρ .

Proof. The results for z_i^* and β^* follow from Lemma B.3. Given $z_i^*(\rho)$, $\omega_i^*(\rho)$ is the unique number between $\underline{x}_i$ and $\bar{x}_i$ satisfying $\mathbb{E}_{\theta_i}[\Psi_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i^*(\rho)) \mid \underline{x}_i \leq \theta_i \leq \bar{x}_i] = z_i^*(\rho)$, where $\underline{x}_i = \min\{\theta_i \mid \Psi_{i, \frac{w_i}{\rho}}^S(\theta_i) \geq z_i^*(\rho)\}$ and $\bar{x}_i = \max\{\theta_i \mid \Psi_{i, \frac{w_i}{\rho}}^B(\theta_i) \leq z_i^*(\rho)\}$. Thus, continuity of $z_i^*(\rho)$ implies continuity of $\omega_i^*(\rho)$. $\square$

We now present the proof of Lemma B.3:

Proof of Lemma B.3. If $\omega_i^* \in (\underline{\theta}_i, \bar{\theta}_i)$, $\frac{w_i}{\rho} < 1$, and $z_1^* \neq z_2^*$, then (i) $\omega_i^* \in (\Psi_{i, \frac{w_i}{\rho}}^{S^{-1}}(z_i^*), \Psi_{i, \frac{w_i}{\rho}}^{B^{-1}}(z_i^*))$, (ii) for all $\theta_i \in (\Psi_{i, \frac{w_i}{\rho}}^{S^{-1}}(z_i^*), \Psi_{i, \frac{w_i}{\rho}}^{B^{-1}}(z_i^*))$,

$$q_i(\theta_i; \rho) = \begin{cases} F_j(\Psi_{j, \frac{w_j}{\rho}}^{S^{-1}}(z_i^*)) & \text{if } z_i^* < z_j^*, \\ F_j(\Psi_{j, \frac{w_j}{\rho}}^{B^{-1}}(z_i^*)) & \text{if } z_i^* > z_j^*, \end{cases}$$

where $i \neq j$, and (iii) $q_i(\omega_i^*) = r_i$, which together imply that

$$z_i^* = \begin{cases} \Psi_{j, \frac{w_j}{\rho}}^S(F_j^{-1}(r_i)) & \text{if } z_i^* < z_j^*, \\ \Psi_{j, \frac{w_j}{\rho}}^B(F_j^{-1}(r_i)) & \text{if } z_i^* > z_j^*. \end{cases}$$

If $z_1^* = z_2^* = z^*$, then z^* and β are defined by

$$q_1(\omega_1^*; \rho) = (1 - \beta)F_2\left(\Psi_{2, \frac{w_2}{\rho}}^{S^{-1}}(z^*)\right) + \beta F_2\left(\Psi_{2, \frac{w_2}{\rho}}^{B^{-1}}(z^*)\right) = r,$$

and

$$q_2(\omega_2^*; \rho) = (1 - \beta)F_1\left(\Psi_{1, \frac{w_1}{\rho}}^{B^{-1}}(z^*)\right) + \beta F_1\left(\Psi_{1, \frac{w_1}{\rho}}^{S^{-1}}(z^*)\right) = 1 - r.$$

Below we show that solutions exist for all r in the relevant range, with the implication that both z^* and β are continuous in ρ .

The pair of equations above means that, when they have a solution, z^* is the solution to

$$r = \frac{F_2\left(\Psi_{2, \frac{w_2}{\rho}}^{S^{-1}}(z^*)\right)\left(1 - F_1\left(\Psi_{1, \frac{w_1}{\rho}}^{S^{-1}}(z^*)\right)\right) - F_2\left(\Psi_{2, \frac{w_2}{\rho}}^{B^{-1}}(z^*)\right)\left(1 - F_1\left(\Psi_{1, \frac{w_1}{\rho}}^{B^{-1}}(z^*)\right)\right)}{F_2\left(\Psi_{2, \frac{w_2}{\rho}}^{S^{-1}}(z^*)\right) - F_1\left(\Psi_{1, \frac{w_1}{\rho}}^{S^{-1}}(z^*)\right) - F_2\left(\Psi_{2, \frac{w_2}{\rho}}^{B^{-1}}(z^*)\right) + F_1\left(\Psi_{1, \frac{w_1}{\rho}}^{B^{-1}}(z^*)\right)} \quad (\text{B.2})$$

and so

$$\beta = \frac{r - F_2\left(\Psi_{2, \frac{w_2}{\rho}}^{S^{-1}}(z^*)\right)}{F_2\left(\Psi_{2, \frac{w_2}{\rho}}^{B^{-1}}(z^*)\right) - F_2\left(\Psi_{2, \frac{w_2}{\rho}}^{S^{-1}}(z^*)\right)}.$$

Thus, we can define (z_1^*, z_2^*) by

$$(z_1^*, z_2^*) = \begin{cases} \left(\Psi_{2, \frac{w_2}{\rho}}^S(F_2^{-1}(r)), \Psi_{1, \frac{w_1}{\rho}}^B(F_1^{-1}(1 - r))\right) & \text{if } \Psi_{2, \frac{w_2}{\rho}}^S(F_2^{-1}(r)) < \Psi_{1, \frac{w_1}{\rho}}^B(F_1^{-1}(1 - r)), \\ \left(\Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(r)), \Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - r))\right) & \text{if } \Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - r)) < \Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(r)), \\ (z^*, z^*) & \text{otherwise.} \end{cases} \quad (\text{B.3})$$

Recalling that we assume overlapping supports, i.e., $\underline{\theta}_2 < \bar{\theta}_1$, we first show that (z_1^*, z_2^*) is continuous in r . Start with $r = 0$. Then

$$\Psi_{2, \frac{w_2}{\rho}}^S(F_2^{-1}(r)) = \Psi_{2, \frac{w_2}{\rho}}^S(\underline{\theta}_2) = \underline{\theta}_2 < \bar{\theta}_1 = \Psi_{1, \frac{w_1}{\rho}}^B(\bar{\theta}_2) = \Psi_{1, \frac{w_1}{\rho}}^B(F_1^{-1}(1 - r)),$$

so, by (B.3), we have $(z_1^*, z_2^*) = \left(\Psi_{2, \frac{w_2}{\rho}}^S(F_2^{-1}(r)), \Psi_{1, \frac{w_1}{\rho}}^B(F_1^{-1}(1 - r))\right) = (\underline{\theta}_2, \bar{\theta}_1)$. As we increase r , z_1^* increases and z_2^* decreases until we get to $r = 1$ or we get to $\hat{r} \in (0, 1)$ such that

$$\Psi_{2, \frac{w_2}{\rho}}^S(F_2^{-1}(\hat{r})) = \Psi_{1, \frac{w_1}{\rho}}^B(F_1^{-1}(1 - \hat{r})).$$

In this case, $\Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - \hat{r})) > \Psi_{1, \frac{w_1}{\rho}}^B(F_1^{-1}(1 - \hat{r})) = \Psi_{2, \frac{w_2}{\rho}}^S(F_2^{-1}(\hat{r})) > \Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(\hat{r}))$. Thus, as r increases further, we have by (B.3) that $(z_1^*, z_2^*) = (z^*, z^*)$. To see that this is a continuous change, note that $\Psi_{2, \frac{w_2}{\rho}}^S(F_2^{-1}(\hat{r})) = \Psi_{1, \frac{w_1}{\rho}}^B(F_1^{-1}(1 - \hat{r}))$, and so (B.2) is satisfied at $\hat{r}$ and $z^* = \Psi_{2, \frac{w_2}{\rho}}^S(F_2^{-1}(\hat{r})) = \Psi_{1, \frac{w_1}{\rho}}^B(F_1^{-1}(1 - \hat{r}))$, i.e., $\lim_{r \downarrow \hat{r}} z^* = z_1^*(\hat{r}) = z_2^*(\hat{r})$.

As r continues to increase above $\hat{r}$, $\Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - r))$ decreases and $\Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(r))$ increases. We continue to have $(z_1^*, z_2^*) = (z^*, z^*)$ until we get to $r = 1$ or we get to $\tilde{r} \in (\hat{r}, 1)$

such that

$$\Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - \tilde{r})) = \Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(\tilde{r})).$$

Then, for $r \in (\tilde{r}, 1]$, we have $\Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - r)) < \Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(r))$ and so

$$(z_1^*, z_2^*) = \left(\Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(r)), \Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - r)) \right).$$

To see that this is a continuous change, note that at $r = \tilde{r}$, we have $\Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - \tilde{r})) = \Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(\tilde{r}))$, and so (B.2) is satisfied at $\tilde{r}$ and $z^* = \Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - \tilde{r})) = \Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(\tilde{r}))$, i.e., $\lim_{r \uparrow \tilde{r}} z^* = z_1^*(\tilde{r}) = z_2^*(\tilde{r})$.

As r continues to increase above $\tilde{r}$, we continue to have $\Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - r)) < \Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(r))$, and so $(z_1^*, z_2^*) = \left(\Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(r)), \Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - r)) \right)$. This completes the proof that (z_1^*, z_2^*) is continuous in r . It also demonstrates that (B.2) has a solution for all r in the relevant range, i.e., for all $r \in (\min\{\hat{r}, 1\}, \min\{\tilde{r}, 1\})$.

Given this, we can rewrite (B.3) as in the statement of the lemma, which completes the proof. $\square$

Lemma B.3 extends to the case of nonoverlapping supports with the following adjustments by analogous arguments:

Lemma B.4. *For $n = 2$ and monotone virtual cost and virtual value functions, if $\bar{\theta}_1 \leq \underline{\theta}_2$ and $\Psi_{1, \frac{w_1}{\rho}}^S(\bar{\theta}_1) > \Psi_{2, \frac{w_2}{\rho}}^B(\underline{\theta}_2)$, then*

$$(z_1^*(\rho), z_2^*(\rho)) = \begin{cases} (z^*(\rho), z^*(\rho)) & \text{if } r \in [0, \tilde{r}(\rho)], \\ \left(\Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(r)), \Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - r)) \right) & \text{if } r \in (\tilde{r}(\rho), 1], \end{cases}$$

and for $\bar{\theta}_1 \leq \underline{\theta}_2$ and $\Psi_{1, \frac{w_1}{\rho}}^S(\bar{\theta}_1) \leq \Psi_{2, \frac{w_2}{\rho}}^B(\underline{\theta}_2)$, we have for all $r \in [0, 1]$,

$$(z_1^*(\rho), z_2^*(\rho)) = \left(\Psi_{2, \frac{w_2}{\rho}}^B(F_2^{-1}(r)), \Psi_{1, \frac{w_1}{\rho}}^S(F_1^{-1}(1 - r)) \right).$$

Combining Lemmas B.3 and B.4, we have shown that z_i^* is continuous in ρ , and so $\omega_i^*(\rho)$ and $\beta^*(\rho)$ are continuous in ρ . The expected budget surplus under binding IR for agents' worst-off types is $\sum_{i \in \mathcal{N}} \Pi_i(\rho)$, where

$$\Pi_i(\rho) \equiv \iint_{\Theta} Q_i(\theta; \rho) \Psi_i(\theta_i, \omega_i^*(\rho)) dF_1(\theta_1) dF_2(\theta_2) - r_i \omega_i^*(\rho),$$

where

$$\mathcal{A}_i^1(\rho) \equiv \{\theta \in \Theta \mid \bar{\Psi}_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i^*(\rho)) > \bar{\Psi}_{j, \frac{w_j}{\rho}}(\theta_j, \omega_j^*(\rho))\},$$

$$\mathcal{A}_i^0(\rho) \equiv \{\theta \in \Theta \mid \bar{\Psi}_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i^*(\rho)) < \bar{\Psi}_{j, \frac{w_j}{\rho}}(\theta_j, \omega_j^*(\rho))\},$$

$$\mathcal{A}_i^\beta(\rho) \equiv \{\boldsymbol{\theta} \in \Theta \mid \bar{\Psi}_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i^*(\rho)) = \bar{\Psi}_{j, \frac{w_j}{\rho}}(\theta_j, \omega_j^*(\rho))\},$$

and

$$Q_1(\boldsymbol{\theta}; \rho) \equiv \begin{cases} 0 & \text{if } \boldsymbol{\theta} \in \mathcal{A}_1^0(\rho), \\ 0 & \text{if } \boldsymbol{\theta} \in \mathcal{A}_1^\beta(\rho) \text{ and } r = \hat{r}(\rho), \\ \beta^*(\rho) \in (0, 1) & \text{if } \boldsymbol{\theta} \in \mathcal{A}_1^\beta(\rho) \text{ and } r \in (\hat{r}(\rho), \tilde{r}(\rho)), \\ 1 & \text{if } \boldsymbol{\theta} \in \mathcal{A}_1^\beta(\rho) \text{ and } r = \tilde{r}(\rho), \\ 1 & \text{if } \boldsymbol{\theta} \in \mathcal{A}_1^1(\rho), \end{cases}$$

where $\beta^*(\rho)$ is given by (B.1) and $Q_2(\boldsymbol{\theta}; \rho) = 1 - Q_1(\boldsymbol{\theta}; \rho)$.

Focusing on agent 1 (analogous arguments hold for agent 2), we have three cases:

1. $r \leq \hat{r}(\rho)$ or $r > \tilde{r}(\rho)$: If $r < \hat{r}(\rho)$ or $r > \tilde{r}(\rho)$, then $\mathcal{A}_1^\beta(\rho) = \emptyset$; if $r = \hat{r}(\rho)$, then $\beta^*(\rho) = 0$. Thus, we have

$$\begin{aligned} \Pi_1(\rho) &= \iint_{\mathcal{A}_1^1(\rho)} \Psi_1(\theta_1, \omega_1^*(\rho)) dF_1(\theta_1) dF_2(\theta_2) - r_1 \omega_1^*(\rho) \\ &= \iint_{\mathcal{A}_1^1(\rho) \text{ s.t. } \theta_1 < \omega_1^*(\rho)} \Psi_1^S(\theta_1) dF_1(\theta_1) dF_2(\theta_2) + \iint_{\mathcal{A}_1^1(\rho) \text{ s.t. } \theta_1 > \omega_1^*(\rho)} \Psi_1^B(\theta_1) dF_1(\theta_1) dF_2(\theta_2) - r_1 \omega_1^*(\rho), \end{aligned}$$

where $\Psi_1^S(\theta_1) = \theta_1 + \frac{F_1(\theta_1)}{f_1(\theta_1)}$ and $\Psi_2^S(\theta_1) = \theta_1 - \frac{1-F_1(\theta_1)}{f_1(\theta_1)}$.

2. $r \in (\hat{r}(\rho), \tilde{r}(\rho))$:

$$\begin{aligned} \Pi_1(\rho) &= \iint_{\mathcal{A}_1^1(\rho)} \Psi_1(\theta_1, \omega_1^*(\rho)) dF_1(\theta_1) dF_2(\theta_2) + \beta^*(\rho) \iint_{\mathcal{A}_1^\beta(\rho)} \Psi_1(\theta_1, \omega_1^*(\rho)) dF_1(\theta_1) dF_2(\theta_2) - r_1 \omega_1^*(\rho) \\ &= \iint_{\mathcal{A}_1^1(\rho) \text{ s.t. } \theta_1 < \omega_1^*(\rho)} \Psi_1^S(\theta_1) dF_1(\theta_1) dF_2(\theta_2) + \iint_{\mathcal{A}_1^1(\rho) \text{ s.t. } \theta_1 > \omega_1^*(\rho)} \Psi_1^B(\theta_1) dF_1(\theta_1) dF_2(\theta_2) \\ &\quad + \beta^*(\rho) \left(\iint_{\mathcal{A}_1^\beta(\rho) \text{ s.t. } \theta_1 < \omega_1^*(\rho)} \Psi_1^S(\theta_1) dF_1(\theta_1) dF_2(\theta_2) + \iint_{\mathcal{A}_1^\beta(\rho) \text{ s.t. } \theta_1 > \omega_1^*(\rho)} \Psi_1^B(\theta_1) dF_1(\theta_1) dF_2(\theta_2) \right) \\ &\quad - r_1 \omega_1^*(\rho). \end{aligned}$$

3. $r = \tilde{r}(\rho)$: In this case, $\beta^*(\rho) = 1$, and so we have

$$\begin{aligned} \Pi_1(\rho) &= \iint_{\mathcal{A}_1^1(\rho) \cup \mathcal{A}_1^\beta(\rho)} \Psi_1(\theta_1, \omega_1^*(\rho)) dF_1(\theta_1) dF_2(\theta_2) - r_1 \omega_1^*(\rho) \\ &= \iint_{\mathcal{A}_1^1(\rho) \cup \mathcal{A}_1^\beta(\rho) \text{ s.t. } \theta_1 < \omega_1^*(\rho)} \Psi_1^S(\theta_1) dF_1(\theta_1) dF_2(\theta_2) \\ &\quad + \iint_{\mathcal{A}_1^1(\rho) \cup \mathcal{A}_1^\beta(\rho) \text{ s.t. } \theta_1 > \omega_1^*(\rho)} \Psi_1^B(\theta_1) dF_1(\theta_1) dF_2(\theta_2) - r_1 \omega_1^*(\rho). \end{aligned}$$

For ρ_0 such that $r < \hat{r}(\rho_0)$, $r > \tilde{r}(\rho_0)$, or $r \in (\hat{r}(\rho_0), \tilde{r}(\rho_0))$, the argument for continuity is straightforward because for a sufficiently small neighborhood of ρ_0 , the applicable case above (either 1 or 2) remains unchanged for any sequence $(\rho_k) \rightarrow \rho$ in that neighborhood, and so we have $\lim_{\rho \rightarrow \rho_0} \Pi(\rho) = \Pi(\rho_0)$.

For ρ_0 such that $r = \hat{r}(\rho_0)$, for all $\rho < \rho_0$, we have $r < \hat{r}(\rho)$, and so again the applicable case above remains unchanged (case 1). For ρ close to but greater than ρ_0 , we have $r \in (\hat{r}(\rho), \tilde{r}(\rho))$, so at the limit, we switch from case 2 to case 1, but because β^* is continuous with $\beta^*(\rho_0) = 0$, we still have $\lim_{\rho \rightarrow \rho_0} \Pi(\rho) = \Pi(\rho_0)$.

For ρ_0 such that $r = \tilde{r}(\rho_0)$, for all ρ sufficiently close to but greater than ρ_0 , we have $r \in (\hat{r}(\rho), \tilde{r}(\rho))$ and so case 2 applies along the sequence and case 3 applies at the limit. Because $\beta^*(\rho_0) = 1$, we have $\lim_{\rho \rightarrow \rho_0} \Pi(\rho) = \Pi(\rho_0)$. Finally, consider the case of $r = \tilde{r}(\rho_0)$, but with a sequence with (ρ_k) converging to ρ_0 from below. Then we have $r > \tilde{r}(\rho_k)$, and so case 1 applies along the sequence and case 3 applies at the limit. At the limit, as ρ increases to ρ_0 , $\mathcal{A}_1^1(\rho)$ contracts by exactly $\mathcal{A}_1^\beta(\rho)$ (i.e., the box shown in Figure B.1(c) appears). More formally, $\lim_{\rho \uparrow \rho_0} \mathcal{A}_1^1(\rho) = \mathcal{A}_1^1(\rho_0) \cup \mathcal{A}_1^\beta(\rho_0)$. And, thus, $\lim_{\rho \rightarrow \rho_0} \Pi(\rho) = \Pi(\rho_0)$.

This completes the proof that $\Pi_i(\rho)$ is continuous. ■

B.3.3 Generalizations

We now discuss how and why the argument generalizes to the case with $n > 2$ and with nonmonotone virtual types, beginning with the former. The piecewise virtual type functions $\Psi_i(\theta_i, \omega_i^*(\rho))$ are not affected by the number of agents other than via the way sets of agents, their weights, shares, and distributions affect the ironing parameters $z_i^*(\rho)$, for $i \in \{1, \dots, n\}$. So if these ironing parameters remain continuous, then the $\omega_i^*(\rho)$'s remain continuous. If $z_i^*(\rho) \neq z_j^*$ for all $j \neq i$, continuity of $z_i^*(\rho)$ follows by the same basic arguments as in case with $n = 2$. The case with ties is slightly more complicated or cumbersome because with $n \geq 3$, there can be ties with 3 or more agents, so the mechanism needs to spell out the tie-breaking probabilities for any subset of agents that can tie with positive probability. However, the tie-breaking probabilities $\beta_i(\rho)$ remain continuous in ρ . In particular if as $\hat{\rho}$ approaches ρ and $z_j^*(\rho)$ approaches from above the ironing parameter $z^*(\rho)$ that is common to a subset of agents excluding j , $\beta_j(\rho) = 1$ will be the case, and similarly, if $z_j^*(\rho)$ approaches $z^*(\rho)$ from below, $\beta_j(\rho) = 0$ will hold.

We now turn to the generalization to nonmonotone virtual type functions. Observe first that any ties that occur not because of countervailing incentives, i.e., because of ironing over intervals that do not include $\omega_i^*(\rho)$, can be broken arbitrarily, e.g., according to the order the agents are labeled or uniform randomly, just like in Myerson (1981). Second, if one replaces all $\Psi_{i, \frac{w_i}{\rho}}^K(\theta_i)$, where $K \in \{S, B\}$, with their ironed counterparts $\overline{\Psi}_{i, \frac{w_i}{\rho}}^K(\theta_i)$, the proof goes

through as is, with $\bar{\Psi}_{i, \frac{w_i}{\rho}}^{K^{-1}}(z_i^*(\rho)) = \Psi_{i, \frac{w_i}{\rho}}^{K^{-1}}(z_i^*(\rho))$, where $\Psi_{i, \frac{w_i}{\rho}}^{K^{-1}}(z_i^*(\rho))$ is well defined insofar as locally where $\Psi_{i, \frac{w_i}{\rho}}^K(\theta_i) = z_i^*(\rho)$ holds, $\Psi_{i, \frac{w_i}{\rho}}^K(\theta_i)$ is increasing.

Last, we discuss how the arguments extend to the case of decreasing marginal values. In this case, and focusing on the two agent case, the allocation rule is $Q_1(\boldsymbol{\theta}; \rho) = \min\{1, \max\{0, \frac{1}{2}(1 + \bar{\Psi}_{1, \frac{w_1}{\rho}}(\theta_1, \omega_1^*(\rho)) - \bar{\Psi}_{2, \frac{w_2}{\rho}}(\theta_2, \omega_2^*(\rho))\}\}$ and $Q_2(\boldsymbol{\theta}; \rho) = 1 - Q_1(\boldsymbol{\theta}; \rho)$. Because this allocation rule is continuous in ρ , the argument that establishes continuity of the expected revenue in ρ is simpler than in the case with constant marginal values as there is no need to separate the type spaces into different regions over which the allocation rule is continuous.

B.3.4 Illustrations

We now briefly illustrate key properties and workings of the IIB mechanism for the uniform-uniform case and then discuss which of these features are more specific and which ones generalize.

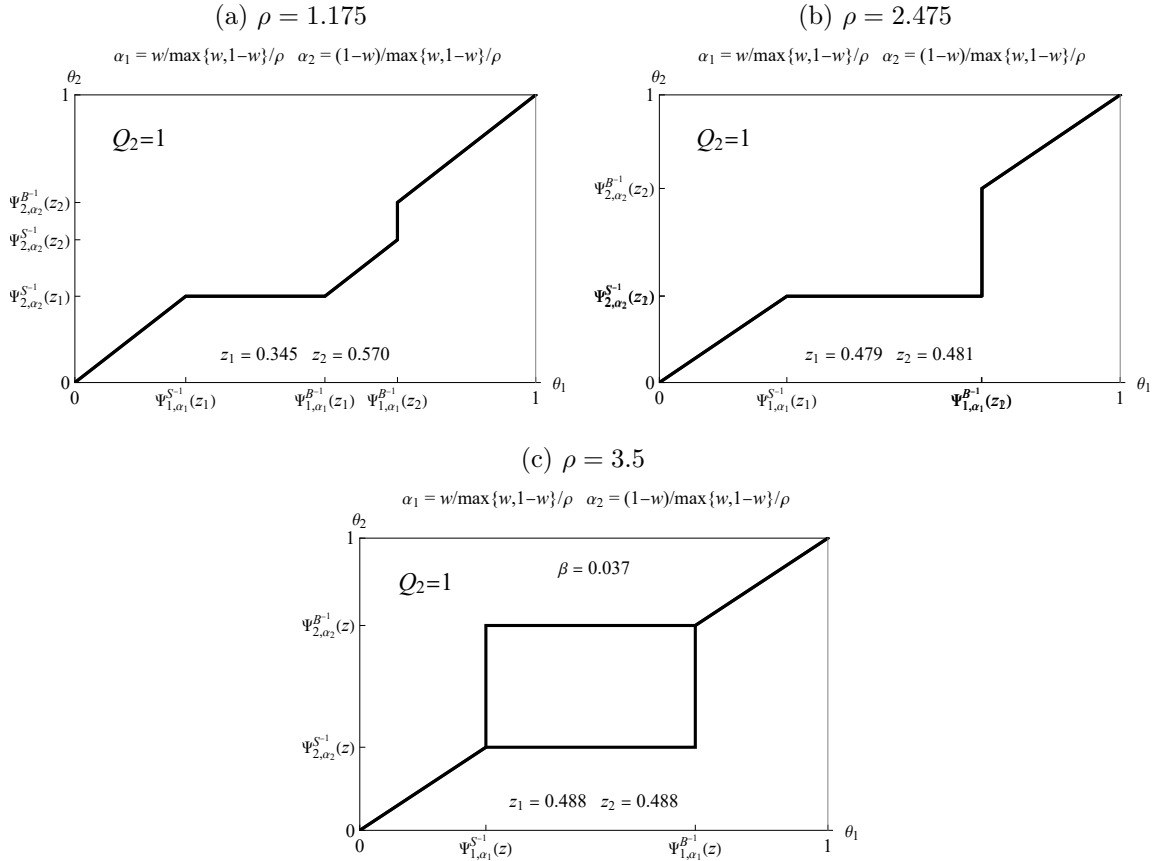


Figure B.1: Assumes $r = 0.3$, $w = 0.4$, and uniformly distributed types on $[0, 1]$.

Figure B.1 displays the ex post allocation for the three different values ρ with $r < 1/2$.

Because of IID, this means that away from the case with ties (displayed in (c)), the boundary line that separates $Q_2 = 1$ from $Q_2 = 0$ has, as θ_i increase, first a flat part before it contains a vertical segment. The height at which the flat part occurs, i.e., the value of θ_2 , does not vary with ρ provided there are no ties because it is given as $F_2^{-1}(r)$. Likewise, the value of θ_1 at which the vertical segment occurs is independent of ρ absent ties and given as $F_1^{-1}(1 - r)$. As ρ increases, the two ironing intervals expand and eventually, at some value of ρ (2.475 here), $F_2^{-1}(r)$ becomes the lower bound of 2's ironing interval and $F_1^{-1}(1 - r)$ the upper bound of 1's. At this point, the ironing parameters have become the same, so the allocation rule now effectively contains a rectangular region, but at this point all ties are broken in favor of agent 2, i.e., $\beta(2.475) = 0$.

As ρ increases further, both agents continue to have the same ironing parameter (which nevertheless varies with ρ , and agent 1 now also wins ties with positive probability, as displayed in panel (c).

In general, and assuming monotone virtual types, the allocation rule may *never* contain a rectangular region (i.e., a “box”) if $F_2^{-1}(r)$ and $F_1^{-1}(1 - r)$ are sufficiently different and, unless $F_2^{-1}(r) = F_1^{-1}(1 - r)$, the allocation rule does not involve a box (ties) if ρ is sufficiently small and the bargaining weights are the same. Increasing ρ may first lead to the emergence of a box and then further increases to its disappearance.

The general condition for there to never be a box is, for $F_2^{-1}(r) < F_1^{-1}(1 - r)$ and assuming that Ψ_1^B and Ψ_2^S are monotone,

$$\Psi_1^B(F_1^{-1}(1 - r)) > F_2^{-1}(r)$$

or

$$\Psi_2^S(F_2^{-1}(r)) < F_1^{-1}(1 - r);$$

for $F_2^{-1}(r) > F_1^{-1}(1 - r)$ and assuming that Ψ_1^S and Ψ_2^B are monotone,

$$\Psi_1^S(F_1^{-1}(1 - r)) < F_2^{-1}(r)$$

or

$$\Psi_2^B(F_2^{-1}(r)) > F_1^{-1}(1 - r).$$

For $F_2^{-1}(r) = F_1^{-1}(1 - r)$, there is always a box as soon as $\rho \geq \max \mathbf{w}$.

B.3.5 Expected revenue is nondecreasing in ρ

Here we prove that $\Pi(\rho)$ is nondecreasing in ρ .

Lemma B.5. $\Pi(\rho)$ is nondecreasing in ρ .

Proof. Recall that the Lagrangian is

$$\mathcal{L} = \sum_{i \in \mathcal{N}} w_i \mathbb{E}_{\boldsymbol{\theta}} [V(\theta_i, Q_i(\boldsymbol{\theta})) - M_i(\boldsymbol{\theta})] + \rho \sum_{i \in \mathcal{N}} \mathbb{E}_{\boldsymbol{\theta}} [M_i(\boldsymbol{\theta})] + \sum_{i \in \mathcal{N}} \gamma_i u_i(\omega_i).$$

To apply the envelope theorem, define

$$V(\rho) \equiv \max_{\mathbf{Q}, \mathbf{M}} \mathcal{L}(\mathbf{Q}, \mathbf{M}; \rho).$$

For fixed $(\mathbf{Q}, \mathbf{M})$, $\mathcal{L}$ is affine in ρ , and so $V(\rho)$ is the pointwise supremum of affine functions and hence convex, which implies that $V'(\rho)$ exists almost everywhere and is nondecreasing. By the envelope theorem, $V'(\rho) = \Pi(\rho)$, and so Π is nondecreasing in ρ .

Alternatively, one can obtain this result as follows: If $(\mathbf{Q}_{\rho}^*, \mathbf{M}_{\rho}^*)$ is the IIB mechanism for a given ρ , then for some function h ,

$$(\mathbf{Q}_{\rho}^*, \mathbf{M}_{\rho}^*) \in \arg \max_{\mathbf{Q}, \mathbf{M}} h(\boldsymbol{\theta}, \mathbf{Q}(\boldsymbol{\theta}), \mathbf{M}(\boldsymbol{\theta})) + \rho \sum_{i \in \mathcal{N}} \mathbb{E}_{\boldsymbol{\theta}} [M_i(\boldsymbol{\theta})].$$

Fixing $\rho_1 < \rho_2$, we have:

$$h(\boldsymbol{\theta}, \mathbf{Q}_{\rho_1}^*(\boldsymbol{\theta}), \mathbf{M}_{\rho_1}^*(\boldsymbol{\theta})) + \rho_1 \sum_{i \in \mathcal{N}} \mathbb{E}_{\boldsymbol{\theta}} [M_{i, \rho_1}^*(\boldsymbol{\theta})] \geq h(\boldsymbol{\theta}, \mathbf{Q}_{\rho_2}^*(\boldsymbol{\theta}), \mathbf{M}_{\rho_2}^*(\boldsymbol{\theta})) + \rho_1 \sum_{i \in \mathcal{N}} \mathbb{E}_{\boldsymbol{\theta}} [M_{i, \rho_2}^*(\boldsymbol{\theta})]$$

and

$$h(\boldsymbol{\theta}, \mathbf{Q}_{\rho_2}^*(\boldsymbol{\theta}), \mathbf{M}_{\rho_2}^*(\boldsymbol{\theta})) + \rho_2 \sum_{i \in \mathcal{N}} \mathbb{E}_{\boldsymbol{\theta}} [M_{i, \rho_2}^*(\boldsymbol{\theta})] \geq h(\boldsymbol{\theta}, \mathbf{Q}_{\rho_1}^*(\boldsymbol{\theta}), \mathbf{M}_{\rho_1}^*(\boldsymbol{\theta})) + \rho_2 \sum_{i \in \mathcal{N}} \mathbb{E}_{\boldsymbol{\theta}} [M_{i, \rho_1}^*(\boldsymbol{\theta})].$$

Subtracting, we get

$$(\rho_2 - \rho_1) \left(\sum_{i \in \mathcal{N}} \mathbb{E}_{\boldsymbol{\theta}} [M_{i, \rho_2}^*(\boldsymbol{\theta})] - \sum_{i \in \mathcal{N}} \mathbb{E}_{\boldsymbol{\theta}} [M_{i, \rho_1}^*(\boldsymbol{\theta})] \right) \geq 0,$$

which implies that $\Pi(\rho)$ is nondecreasing. ■

B.4 Proof of Proposition 1

Proof of Proposition 1. The first part of the proposition follows from Theorem 1. It remains to show that if $n = 2$ and virtual values and virtual costs are increasing, then ties between agents' ironed weighted virtual types occur with probability zero. We begin by noting the following two helpful relations: for $p \in [\underline{\theta}_i, \bar{\theta}_i]$,

$$\int_{\underline{\theta}_i}^p \Psi_i^S(\theta) dF_i(\theta) = p F_i(p) \tag{B.4}$$

and

$$\int_p^{\bar{\theta}_i} \Psi_i^B(\theta) dF_i(\theta) = p(1 - F_i(p)). \tag{B.5}$$

We assume that the virtual cost and virtual value functions, Ψ_i^S and Ψ_i^B , are monotone. We generalize the assumption on the supports of the agents' distributions to allow any bounded interval supports $[\underline{\theta}_1, \bar{\theta}_1]$ and $[\underline{\theta}_2, \bar{\theta}_2]$, not just $[\underline{\theta}_1, \bar{\theta}_1] = [0, 1]$ and $[\underline{\theta}_2, \bar{\theta}_2] = [\underline{\theta}, 1 + \underline{\theta}]$ as assumed in the body of the paper.

Assume $n = 2$. Wlog, assume that $0 \leq w_2 \leq w_1 = 1$. That is, assume that agent 1 has the weakly higher bargaining weight and normalize the weights so that the largest weight is equal to 1. We let ρ denote the Lagrange multiplier on the no-deficit constraint. Because $\rho \geq \max\{w_1, w_2\}$ is required for feasibility, under our normalization, we have $\rho \geq 1$. Letting $\Pi_\rho(r_1, w_2)$ denote the expected budget surplus under binding IR for the agents' worst-off types given ρ , r_1 , w_1 , and w_2 , we have

$$\Pi_\rho(r_1, w_2) = \sum_{i=1}^2 \int_{\underline{\theta}_i}^{\bar{\theta}_i} \Psi_i(\theta, \hat{\theta}_i) q_i(\theta) dF_i(\theta) - r_1 \hat{\theta}_1 - (1 - r_1) \hat{\theta}_2,$$

where $\hat{\theta}_i$ is agent i 's worst-off type and where $\hat{\theta}_1$, $\hat{\theta}_2$, q_i , and q_2 depend on ρ , r_1 , w_1 , and w_2 . Define $\alpha_1 \equiv w_1/\rho = 1/\rho$ and $\alpha_2 \equiv w_2/\rho$.

Suppose, by way of contradiction, that ties occur with positive probability. Then it must be that both agents' weighted virtual type functions require ironing and have a common ironing parameter, $z_1 = z_2 = z$, where $z \in [\max\{\Psi_{1,\alpha_1}^B(\underline{\theta}_1), \Psi_{2,\alpha_2}^B(\underline{\theta}_2)\}, \min\{\Psi_{1,\alpha_1}^S(\bar{\theta}_1), \Psi_{2,\alpha_2}^S(\bar{\theta}_2)\}]$. Further, we must have $\rho > 1$ because if $\rho = 1$, then agent 1's weighted virtual type functions are the identity functions, $\Psi_{1,1/\rho}^S(\theta) = \Psi_{1,1/\rho}^B(\theta) = \Psi_{1,1/\rho}(\theta; \hat{\theta}_1) = \theta$, and so no ironing is required for agent 1. The requirement that $\rho > 1$ means that we must have $\Pi_1(r_1, w_2) < 0$ because ρ is the smallest multiplier greater than or equal to $\max\{w_1, w_2\}$ such that $\Pi_\rho(r_1, w_2) \geq 0$.

In what follows, it will be useful to define the inverse weighted virtual type functions to have a domain of the entire real line. To do this, we abuse notation by defining for $J \in \{S, B\}$ and $i \in \{1, 2\}$,

$$\Psi_{i,\alpha_i}^{J^{-1}}(x) \equiv \begin{cases} \underline{\theta}_i & \text{if } x < \Psi_{i,\alpha_i}^J(\underline{\theta}_i), \\ \Psi_{i,\alpha_i}^{J^{-1}}(x) & \text{if } x \in [\Psi_{i,\alpha_i}^J(\underline{\theta}_i), \Psi_{i,\alpha_i}^J(\bar{\theta}_i)], \\ \bar{\theta}_i & \text{if } x > \Psi_{i,\alpha_i}^J(\bar{\theta}_i). \end{cases}$$

Let $a \in [0, 1]$ denote the probability with which agent 1 wins a tie. Then the agents' worst-off types satisfy

$$q_1(\hat{\theta}_1) = r_1 = F_2(\Psi_{2,\alpha_2}^{S^{-1}}(z)) + a \left(F_2(\Psi_{2,\alpha_2}^{B^{-1}}(z)) - F_2(\Psi_{2,\alpha_2}^{S^{-1}}(z)) \right)$$

and

$$q_2(\hat{\theta}_2) = 1 - r_1 = F_1(\Psi_{1,\alpha_1}^{S^{-1}}(z)) + (1 - a) \left(F_1(\Psi_{1,\alpha_1}^{B^{-1}}(z)) - F_1(\Psi_{1,\alpha_1}^{S^{-1}}(z)) \right).$$

Further, q_i is constant in the ironing interval: for all $\theta_i \in [\Psi_{i,\alpha_i}^{S^{-1}}(z), \Psi_{i,\alpha_i}^{B^{-1}}(z)]$, $q_i(\theta_i) = q_i(\hat{\theta}_i) = r_i$. Solving the above displayed equations for a , we have $a = \frac{r_1 - (1 - F_1(\Psi_{1,\alpha_1}^{B^{-1}}(z)))}{F_1(\Psi_{1,\alpha_1}^{B^{-1}}(z)) - F_1(\Psi_{1,\alpha_1}^{S^{-1}}(z))}$. Because $a \in [0, 1]$, parameters must be such that

$$1 - F_1(\Psi_{1,\alpha_1}^{B^{-1}}(z)) \leq r_1 \leq 1 - F_1(\Psi_{1,\alpha_1}^{S^{-1}}(z)). \quad (\text{B.6})$$

Under our assumption of a common ironing parameter z , we have for $\theta \in [\underline{\theta}_1, \bar{\theta}_1]$,

$$q_1(\theta) = \begin{cases} F_2\left(\Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(\theta))\right) & \text{if } \theta \in [0, \Psi_{1,\alpha_1}^{S^{-1}}(z)], \\ r_1 & \text{if } \theta \in [\Psi_{1,\alpha_1}^{S^{-1}}(z), \Psi_{1,\alpha_1}^{B^{-1}}(z)], \\ F_2\left(\Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(\theta))\right) & \text{if } \theta \in [\Psi_{1,\alpha_1}^{B^{-1}}(z), 1], \end{cases}$$

and analogously for $q_2(\theta)$.

We now write out the expression for the expected budget surplus $\Pi_\rho(r_1, w_2)$. We do this in two pieces, defining $\Pi_\rho(r_1, w_2) \equiv \Pi_{\rho,1}(r_1, w_2) + \Pi_{\rho,2}(r_1, w_2)$, where $\Pi_{\rho,i}(r_1, w_2) \equiv \int_{\underline{\theta}_i}^{\bar{\theta}_i} q_i(\theta) \Psi_i(\theta, \hat{\theta}_i) dF_i(\theta) - r_i \hat{\theta}_i$. Starting with $\Pi_{\rho,1}$, we have:

$$\begin{aligned} \Pi_{\rho,1}(r_1, w_2) &= \int_{\underline{\theta}_1}^{\Psi_{1,\alpha_1}^{S^{-1}}(z)} \Psi_1^S(\theta) F_2\left(\Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(\theta))\right) dF_1(\theta) + r_1 \int_{\Psi_{1,\alpha_1}^{S^{-1}}(z)}^{\Psi_{1,\alpha_1}^{B^{-1}}(z)} \Psi_1(\theta, \hat{\theta}_1) dF_1(\theta) \\ &\quad + \int_{\Psi_{1,\alpha_1}^{B^{-1}}(z)}^{\bar{\theta}_1} \Psi_1^B(\theta) F_2\left(\Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(\theta))\right) dF_1(\theta) - r_1 \hat{\theta}_1. \end{aligned}$$

The integral over $[\Psi_{1,\alpha_1}^{S^{-1}}(z), \Psi_{1,\alpha_1}^{B^{-1}}(z)]$ can be rewritten as follows:

$$\begin{aligned} &\int_{\Psi_{1,\alpha_1}^{S^{-1}}(z)}^{\Psi_{1,\alpha_1}^{B^{-1}}(z)} \Psi_1(\theta, \hat{\theta}_1) dF_1(\theta) \\ &= \int_{\Psi_{1,\alpha_1}^{S^{-1}}(z)}^{\hat{\theta}_1} \Psi_1^S(\theta) dF_1(\theta) + \int_{\hat{\theta}_1}^{\Psi_{1,\alpha_1}^{B^{-1}}(z)} \Psi_1^B(\theta) dF_1(\theta) \\ &= \int_{\underline{\theta}_1}^{\hat{\theta}_1} \Psi_1^S(\theta) dF_1(\theta) - \int_{\underline{\theta}_1}^{\Psi_{1,\alpha_1}^{S^{-1}}(z)} \Psi_1^S(\theta) dF_1(\theta) + \int_{\hat{\theta}_1}^{\bar{\theta}_1} \Psi_1^B(\theta) dF_1(\theta) - \int_{\Psi_{1,\alpha_1}^{B^{-1}}(z)}^{\bar{\theta}_1} \Psi_1^B(\theta) dF_1(\theta) \\ &= \hat{\theta}_1 F_1(\hat{\theta}_1) - \Psi_{1,\alpha_1}^{S^{-1}}(z) F_1(\Psi_{1,\alpha_1}^{S^{-1}}(z)) + \hat{\theta}_1 (1 - F_1(\hat{\theta}_1)) - \Psi_{1,\alpha_1}^{B^{-1}}(z) (1 - F_1(\Psi_{1,\alpha_1}^{B^{-1}}(z))) \\ &= \hat{\theta}_1 - \Psi_{1,\alpha_1}^{S^{-1}}(z) F_1(\Psi_{1,\alpha_1}^{S^{-1}}(z)) - \Psi_{1,\alpha_1}^{B^{-1}}(z) (1 - F_1(\Psi_{1,\alpha_1}^{B^{-1}}(z))) \\ &= \hat{\theta}_1 - X_1, \end{aligned}$$

where the third equality uses (B.4) and (B.5) and where

$$X_i \equiv \Psi_{i,\alpha_i}^{S^{-1}}(z) F_i(\Psi_{i,\alpha_i}^{S^{-1}}(z)) + \Psi_{i,\alpha_i}^{B^{-1}}(z) (1 - F_i(\Psi_{i,\alpha_i}^{B^{-1}}(z))). \quad (\text{B.7})$$

Thus, we have the following expression for $\Pi_{\rho,1}(r_1, w_2)$, which no longer involves $\hat{\theta}_1$:

$$\begin{aligned}\Pi_{\rho,1}(r_1, w_2) &= \int_{\underline{\theta}_1}^{\Psi_{1,\alpha_1}^{S^{-1}}(z)} \Psi_1^S(\theta) F_2\left(\Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(\theta))\right) dF_1(\theta) \\ &\quad - \int_{\Psi_{1,\alpha_1}^{B^{-1}}(z)}^{\bar{\theta}_1} \Psi_1^B(\theta) \left[1 - F_2\left(\Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(\theta))\right)\right] dF_1(\theta) \\ &\quad + \Psi_{1,\alpha_1}^{B^{-1}}(z) \left[1 - F_1(\Psi_{1,\alpha_1}^{B^{-1}}(z))\right] - r_1 X_1.\end{aligned}$$

Repeating this for firm 2, we obtain an analogous expression:

$$\begin{aligned}\Pi_{\rho,2}(r_1, w_2) &= \int_{\underline{\theta}_2}^{\Psi_{2,\alpha_2}^{S^{-1}}(z)} \Psi_2^S(\theta) F_1\left(\Psi_{1,\alpha_1}^{S^{-1}}(\Psi_{2,\alpha_2}^S(\theta))\right) dF_2(\theta) \\ &\quad - \int_{\Psi_{2,\alpha_2}^{B^{-1}}(z)}^{\bar{\theta}_2} \Psi_2^B(\theta) \left[1 - F_1\left(\Psi_{1,\alpha_1}^{B^{-1}}(\Psi_{2,\alpha_2}^B(\theta))\right)\right] dF_2(\theta) \\ &\quad + \Psi_{2,\alpha_2}^{B^{-1}}(z) \left[1 - F_2(\Psi_{2,\alpha_2}^{B^{-1}}(z))\right] - (1 - r_1) X_2.\end{aligned}$$

For each of the two integrals in the above expression for $\Pi_{\rho,2}(r_1, w_2)$, we can perform a change of variables, using $\theta = \Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(y))$ and $\theta = \Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(y))$, respectively, to get

$$\begin{aligned}&\int_{\underline{\theta}_2}^{\Psi_{2,\alpha_2}^{S^{-1}}(z)} \Psi_2^S(\theta) F_1\left(\Psi_{1,\alpha_1}^{S^{-1}}(\Psi_{2,\alpha_2}^S(\theta))\right) dF_2(\theta) \\ &= \Psi_{2,\alpha_2}^{S^{-1}}(z) F_1\left(\Psi_{1,\alpha_1}^{S^{-1}}(z)\right) F_2(\Psi_{2,\alpha_2}^{S^{-1}}(z)) - \int_{\Psi_{1,\alpha_1}^{S^{-1}}(\underline{\theta}_2)}^{\Psi_{1,\alpha_1}^{S^{-1}}(z)} \Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(y)) F_2(\Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(y))) dF_1(y)\end{aligned}$$

and

$$\begin{aligned}&\int_{\Psi_{2,\alpha_2}^{B^{-1}}(z)}^{\bar{\theta}_2} \Psi_2^B(\theta) \left[1 - F_1\left(\Psi_{1,\alpha_1}^{B^{-1}}(\Psi_{2,\alpha_2}^B(\theta))\right)\right] dF_2(\theta) \\ &= -\Psi_{2,\alpha_2}^{B^{-1}}(z) \left[1 - F_1\left(\Psi_{1,\alpha_1}^{B^{-1}}(z)\right)\right] F_2(\Psi_{2,\alpha_2}^{B^{-1}}(z)) \\ &\quad - \int_{\Psi_{1,\alpha_1}^{B^{-1}}(z)}^{\Psi_{1,\alpha_1}^{B^{-1}}(z)(\bar{\theta}_2)} \Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(y)) \left[1 - F_2(\Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(y)))\right] dF_1(y) \\ &\quad + \Psi_{1,\alpha_1}^{B^{-1}}(z) \left[1 - F_1(\Psi_{1,\alpha_1}^{B^{-1}}(z))\right] - \int_{\Psi_{2,\alpha_2}^{B^{-1}}(z)}^{\bar{\theta}_2} \left[1 - F_1\left(\Psi_{1,\alpha_1}^{B^{-1}}(\Psi_{2,\alpha_2}^B(\theta))\right)\right] d\theta.\end{aligned}$$

Making these substitutions, and simplifying, we get

$$\begin{aligned}
\Pi_\rho(r_1, w_2) &= \Pi_{\rho,1}(r_1, w_2) + \Pi_{\rho,2}(r_1, w_2) \\
&= \int_{\max\{\underline{\theta}_1, \Psi_{1,\alpha_1}^{S^{-1}}(\underline{\theta}_2)\}}^{\Psi_{1,\alpha_1}^{S^{-1}}(z)} \left[\Psi_1^S(\theta) - \Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(\theta)) \right] F_2 \left(\Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(\theta)) \right) dF_1(\theta) \\
&\quad + \int_{\Psi_{1,\alpha_1}^{B^{-1}}(z)}^{\min\{\bar{\theta}_1, \Psi_{1,\alpha_1}^{B^{-1}}(\bar{\theta}_2)\}} \left[\Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(\theta)) - \Psi_1^B(\theta) \right] \left(1 - F_2(\Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(\theta))) \right) dF_1(\theta) \\
&\quad + \Psi_{2,\alpha_2}^{S^{-1}}(z) F_2(\Psi_{2,\alpha_2}^{S^{-1}}(z)) F_1 \left(\Psi_{1,\alpha_1}^{S^{-1}}(z) \right) - \Psi_{2,\alpha_2}^{B^{-1}}(z) F_2(\Psi_{2,\alpha_2}^{B^{-1}}(z)) F_1 \left(\Psi_{1,\alpha_1}^{B^{-1}}(z) \right) \\
&\quad + \Psi_{2,\alpha_2}^{B^{-1}}(z) + \int_{\Psi_{2,\alpha_2}^{B^{-1}}(z)}^{\bar{\theta}_2} \left[1 - F_1 \left(\Psi_{1,\alpha_1}^{B^{-1}}(\Psi_{2,\alpha_2}^B(\theta)) \right) \right] d\theta \\
&\quad - r_1 X_1 - (1 - r_1) X_2.
\end{aligned}$$

The above expression requires some care in dealing with differing supports. If $\underline{\theta}_1 < \underline{\theta}_2$, then $\Psi_{1,\alpha_1}^{S^{-1}}(\underline{\theta}_2) > \underline{\theta}_1$, and if $\underline{\theta}_1 \geq \underline{\theta}_2$, then $\Psi_{1,\alpha_1}^{S^{-1}}(\underline{\theta}_2) = \underline{\theta}_1$, and so the sum of

$$\int_{\underline{\theta}_1}^{\Psi_{1,\alpha_1}^{S^{-1}}(z)} \Psi_1^S(\theta) F_2(\Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(\theta))) dF_1(\theta)$$

and

$$- \int_{\Psi_{1,\alpha_1}^{S^{-1}}(\underline{\theta}_2)}^{\Psi_{1,\alpha_1}^{S^{-1}}(z)} \Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(y)) F_2(\Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(y))) dF_1(y)$$

is

$$\begin{aligned}
&\int_{\underline{\theta}_1}^{\max\{\underline{\theta}_1, \Psi_{1,\alpha_1}^{S^{-1}}(\underline{\theta}_2)\}} \Psi_1^S(\theta) F_2 \left(\Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(\theta)) \right) dF_1(\theta) \\
&\quad + \int_{\max\{\underline{\theta}_1, \Psi_{1,\alpha_1}^{S^{-1}}(\underline{\theta}_2)\}}^{\Psi_{1,\alpha_1}^{S^{-1}}(z)} \left[\Psi_1^S(\theta) - \Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(\theta)) \right] F_2 \left(\Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(\theta)) \right) dF_1(\theta),
\end{aligned}$$

but the integrand in the first integral is zero for all θ within the bounds of integration, leaving us with the expression above in the equation for $\Pi_\rho(r_1, w_2)$. Analogously, if $\bar{\theta}_1 > \bar{\theta}_2$, then $\Psi_{1,\alpha_1}^{B^{-1}}(\bar{\theta}_2) < \bar{\theta}_1$, and if $\bar{\theta}_1 \leq \bar{\theta}_2$, then $\Psi_{1,\alpha_1}^{B^{-1}}(\bar{\theta}_2) = \bar{\theta}_1$, and so the sum of

$$- \int_{\Psi_{1,\alpha_1}^{B^{-1}}(z)}^{\bar{\theta}_1} \Psi_1^B(\theta) [1 - F_2(\Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(\theta)))] dF_1(\theta)$$

and

$$\int_{\Psi_{1,\alpha_1}^{B^{-1}}(z)}^{\Psi_{1,\alpha_1}^{B^{-1}}(\bar{\theta}_2)} \Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(y)) \left[1 - F_2(\Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(y))) \right] dF_1(y)$$

is

$$\begin{aligned}
& - \int_{\min\{\bar{\theta}_1, \Psi_{1,\alpha_1}^{B^{-1}}(\bar{\theta}_2)\}}^{\bar{\theta}_1} \Psi_1^B(\theta) \left[1 - F_2(\Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(\theta))) \right] dF_1(\theta) \\
& + \int_{\Psi_{1,\alpha_1}^{B^{-1}}(z)}^{\min\{\bar{\theta}_1, \Psi_{1,\alpha_1}^{B^{-1}}(\bar{\theta}_2)\}} \left[\Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(\theta)) - \Psi_1^B(\theta) \right] \left[1 - F_2(\Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(\theta))) \right] dF_1(\theta),
\end{aligned}$$

but the integrand in the first integral is zero for all θ in the bounds of integration, and so agent we are left with the expression above in the equation for $\Pi_\rho(r_1, w_2)$.

We now ask the question whether $\Pi_1(r_1, w_2) \geq 0$, which would be a contradiction. Under the assumption that $\rho = 1$, we have $\alpha_1 = 1$ (and $\alpha_2 < 1$) and so virtual types with weight α_1 are the identity function. Given this, $\underline{\theta}_1 < z < \bar{\theta}_1$ and (B.6) implies that $r = 1 - F_1(z)$ and (B.7) implies that $X_1 = z$. Using these and the expression for X_2 from (B.7), we have

$$\begin{aligned}
\Pi_1(r_1, w_2) &= \int_{\max\{\underline{\theta}_1, \underline{\theta}_2\}}^z \underbrace{\left(\Psi_1^S(\theta) - \Psi_{2,\alpha_2}^{S^{-1}}(\theta) \right)}_{\text{positive}} F_2\left(\Psi_{2,\alpha_2}^{S^{-1}}(\theta)\right) dF_1(\theta) \\
&+ \int_z^{\min\{\bar{\theta}_1, \bar{\theta}_2\}} \underbrace{\left(\Psi_{2,\alpha_2}^{B^{-1}}(\theta) - \Psi_1^B(\theta) \right)}_{\text{positive}} \left(1 - F_2(\Psi_{2,\alpha_2}^{B^{-1}}(\theta)) \right) dF_1(\theta) \\
&+ (1 - F_1(z)) \underbrace{\left(\Psi_{2,\alpha_2}^{B^{-1}}(z) - z \right)}_{\text{positive}} + \int_{\Psi_{2,\alpha_2}^{B^{-1}}(z)}^{\bar{\theta}_2} \left(1 - F_1(\Psi_{2,\alpha_2}^B(\theta)) \right) d\theta \\
&\geq 0,
\end{aligned}$$

which contradicts $\rho > 1$. Thus, we conclude that whenever $z_1 = z_2 = z$, we must have $\rho = 1$, which implies that ties are a zero probability event, completing the proof. ■

B.5 Proof of Lemma A.2 (used in the proof of Theorem 2)

Proof of Lemma A.2. Let $\mathcal{N}_U$ denote the set of $n_U \geq 1$ agents with distribution support $[0, 1]$ and $\mathcal{N}_D$ denote the set of $n_D \geq 1$ agents with distribution support $[\underline{\theta}, 1 + \underline{\theta}]$. Define $\mathcal{N} \equiv \mathcal{N}_U \cup \mathcal{N}_D$.

An implication of IC is that $u'_i(\theta) = q_i(\theta) - r_i$ wherever u_i is differentiable, which by IC is almost everywhere. Given this, the monotonicity of u_i implies the following characterization of the set of worst-off types for agent i , denoted by $\Omega_i \equiv \arg \min_{\theta_i \in [\underline{\theta}_i, \bar{\theta}_i]} u_i(\theta_i)$ (see also Cramton et al. (1987, Lemma 2) and Loertscher and Wasser (2019)):

$$\Omega_i = \begin{cases} \{\theta_i \in [\underline{\theta}_i, \bar{\theta}_i] \mid q_i(\theta_i) = r_i\} & \text{if } \exists \theta_i \in [\underline{\theta}_i, \bar{\theta}_i] \text{ s.t. } q_i(\theta_i) = r_i, \\ \{\theta_i \in [\underline{\theta}_i, \bar{\theta}_i] \mid q_i(z) < r_i \ \forall z < \theta_i \text{ and } q_i(z) > r_i \ \forall z > \theta_i\} & \text{otherwise.} \end{cases}$$

In the first case, in which there exists $\theta_i \in [\underline{\theta}_i, \bar{\theta}_i]$ such that $q_i(\theta_i) = r_i$, the set Ω_i is a (possibly degenerate) interval, and in the second case, Ω_i is a singleton.

Using this characterization of worst-off types, in our setting with $\underline{\theta} < 1$, for any agent $j \in \mathcal{N}_U$, $q_j^e(\underline{\theta}) = 0$ and $q_j^e(1) \leq 1$ (strict if $\underline{\theta} > 0$), which implies that either (i) $\hat{\theta}_j^e(r_j) \in [0, 1]$ and $q_j^e(\hat{\theta}_j^e(r_j)) = r_j$ or (ii) $\hat{\theta}_j^e(r_j) = 1$ and $q_j^e(1) < r_j$. If $n_D \geq 2$, then for all $i \in \mathcal{N}_D$, $q_i^e(\underline{\theta}) = 0$ and $q_i^e(1 + \underline{\theta}) = 1$. This implies that for all $i \in \mathcal{N}_D$, $q_i^e(\hat{\theta}_i^e(r_i)) = r_i$. If $n_D = 1$ and $\underline{\theta} > 0$, then for agent $i \in \mathcal{N}_D$, we have $q_i^e(\underline{\theta}) > 0$. This implies that for the sole agent in $\mathcal{N}_D$, either (i) $q_i^e(\hat{\theta}_i^e(r_i)) = r_i$ or (ii) $\hat{\theta}_i^e(r_i) = \underline{\theta}$ and $r_i < q_i^e(\underline{\theta})$.

Given this, the result follows from the combination of Lemmas B.6–B.9, which show that $\Pi_c(\mathbf{r})$ is concave in $\mathbf{r}$ (Lemma B.6), there exists ownership $\mathbf{r}^*$ that equalizes the agents worst-off types (Lemma B.7), $\Pi_c(\mathbf{r})$ is maximized at $\mathbf{r}^*$ (Lemma B.8), and $\Pi_c(\mathbf{r}^*) > 0$ (Lemma B.9). We now state and prove these lemmas.

Lemma B.6. *Given $\underline{\theta} < 1$, $\Pi_c(\mathbf{r})$ is concave in $\mathbf{r}$.*

Proof. Given $\mathbf{r}$, let $\hat{\theta}_i^e(r_i)$ denote a worst-off type of agent i under the ex post efficient allocation rule and ownership $\mathbf{r}$. To show that Π_c is concave, we show that for all $\lambda \in (0, 1)$ and $\mathbf{r}, \hat{\mathbf{r}} \in \Delta$ (dropping the subscript and superscript on Π_c to ease notation):

$$\Pi(\lambda \mathbf{r} + (1 - \lambda) \hat{\mathbf{r}}) \geq \lambda \Pi(\mathbf{r}) + (1 - \lambda) \Pi(\hat{\mathbf{r}}). \quad (\text{B.8})$$

Note that Π can be written as $\Pi(\mathbf{r}) = \sum_{i \in \mathcal{N}} \Pi_i(r_i)$, where

$$\Pi_i(r_i) \equiv \int_{\underline{\theta}_i}^{\hat{\theta}_i^e(r_i)} \Psi_i^S(\theta) q_i^e(\theta) dF_i(\theta) + \int_{\hat{\theta}_i^e(r_i)}^{\bar{\theta}_i} \Psi_i^B(\theta) q_i^e(\theta) dF_i(\theta) - r_i \hat{\theta}_i^e(r_i).$$

Using the definition of $\Pi_i(r_i)$ and $(\Psi_i^S(\theta) - \Psi_i^B(\theta)) f_i(\theta) = 1$, we have

$$\begin{aligned} \Pi_i(\lambda r_i + (1 - \lambda) \hat{r}_i) &= \lambda \Pi_i(r_i) + (1 - \lambda) \Pi_i(\hat{r}_i) \\ &\quad + \lambda \int_{\hat{\theta}_i^e(r_i)}^{\hat{\theta}_i^e(\lambda r_i + (1 - \lambda) \hat{r}_i)} q_i^e(\theta) d\theta - (1 - \lambda) \int_{\hat{\theta}_i^e(\lambda r_i + (1 - \lambda) \hat{r}_i)}^{\hat{\theta}_i^e(\hat{r}_i)} q_i^e(\theta) d\theta \\ &\quad + \lambda r_i \hat{\theta}_i^e(r_i) + (1 - \lambda) \hat{r}_i \hat{\theta}_i^e(\hat{r}_i) - (\lambda r_i + (1 - \lambda) \hat{r}_i) \hat{\theta}_i^e(\lambda r_i + (1 - \lambda) \hat{r}_i). \end{aligned} \quad (\text{B.9})$$

Consider an agent i and let $\lambda \in (0, 1)$ and $r_i, \hat{r}_i \in [0, 1]$ with $0 \leq r_i < \hat{r}_i \leq 1$ be given. There are five possibilities for agent i 's worst-off types: (i) $\hat{\theta}_i^e(r_i) < \hat{\theta}_i^e(\hat{r}_i)$ and $r_i = q_i^e(\hat{\theta}_i^e(r_i)) < q_i^e(\hat{\theta}_i^e(\hat{r}_i)) = \hat{r}_i$; (ii) agent $i \in \mathcal{N}_U$, $\hat{\theta}_i^e(r_i) < \hat{\theta}_i^e(\hat{r}_i) = 1$, and $r_i = q_i^e(\hat{\theta}_i^e(r_i)) < q_i^e(\hat{\theta}_i^e(\hat{r}_i)) < \hat{r}_i$; (iii) agent i is the only agent in $\mathcal{N}_D$, $\underline{\theta} = \hat{\theta}_i^e(r_i) < \hat{\theta}_i^e(\hat{r}_i)$, and $r_i < q_i^e(\underline{\theta}) < q_i^e(\hat{\theta}_i^e(\hat{r}_i)) = \hat{r}_i$; (iv) agent $i \in \mathcal{N}_U$, $\hat{\theta}_i^e(r_i) = \hat{\theta}_i^e(\hat{r}_i) = 1$, and $q_i^e(1) < r_i < \hat{r}_i$; or (v) agent i is the only agent in $\mathcal{N}_D$, $\underline{\theta} = \hat{\theta}_i^e(r_i) = \hat{\theta}_i^e(\hat{r}_i)$, and $r_i < \hat{r}_i < q_i^e(\underline{\theta})$.

Under possibilities (i), (ii), and (iii),

$$\int_{\hat{\theta}_i^e(r_i)}^{\hat{\theta}_i^e(\lambda r_i + (1-\lambda)\hat{r}_i)} q_i^e(\theta) d\theta > r_i \left(\hat{\theta}_i^e(\lambda r_i + (1-\lambda)\hat{r}_i) - \hat{\theta}_i^e(r_i) \right) \quad (\text{B.10})$$

and

$$\int_{\hat{\theta}_i^e(\lambda r_i + (1-\lambda)\hat{r}_i)}^{\hat{\theta}_i^e(\hat{r}_i)} q_i^e(\theta) d\theta \leq \hat{r}_i \left(\hat{\theta}_i^e(\hat{r}_i) - \hat{\theta}_i^e(\lambda r_i + (1-\lambda)\hat{r}_i) \right). \quad (\text{B.11})$$

Thus, under possibilities (i), (ii), and (iii),

$$\begin{aligned} \Pi_i(\lambda r_i + (1-\lambda)\hat{r}_i) &> \lambda \Pi_i(r_i) + (1-\lambda) \Pi_i(\hat{r}_i) \\ &+ \lambda r_i \left(\hat{\theta}_i^e(\lambda r_i + (1-\lambda)\hat{r}_i) - \hat{\theta}_i^e(r_i) \right) - (1-\lambda) \hat{r}_i \left(\hat{\theta}_i^e(\hat{r}_i) - \hat{\theta}_i^e(\lambda r_i + (1-\lambda)\hat{r}_i) \right) \\ &+ \lambda r_i \hat{\theta}_i^e(r_i) + (1-\lambda) \hat{r}_i \hat{\theta}_i^e(\hat{r}_i) - (\lambda r_i + (1-\lambda)\hat{r}_i) \hat{\theta}_i^e(\lambda r_i + (1-\lambda)\hat{r}_i), \end{aligned}$$

which simplifies to

$$\Pi_i(\lambda r_i + (1-\lambda)\hat{r}_i) > \lambda \Pi_i(r_i) + (1-\lambda) \Pi_i(\hat{r}_i).$$

In contrast, under possibilities (iv) and (v),

$$\Pi_i(\lambda r_i + (1-\lambda)\hat{r}_i) = \lambda \Pi_i(r_i) + (1-\lambda) \Pi_i(\hat{r}_i).$$

Thus, $\Pi_i(\lambda r_i + (1-\lambda)\hat{r}_i) \geq \lambda \Pi_i(r_i) + (1-\lambda) \Pi_i(\hat{r}_i)$ for each agent, implying that (B.8) holds and proving that Π is concave. $\square$

Lemma B.7. Assume $\underline{\theta} < 1$, and let $\mathbf{r}^*$ be defined by for all $i \in \mathcal{N}$,

$$r_i^* \equiv \begin{cases} q_i^e(\hat{\theta}) & \text{if } \sum_{i \in \mathcal{N}} q_i^e(1) \geq 1, \\ q_i^e(1) & \text{if } \sum_{i \in \mathcal{N}} q_i^e(1) < 1 \text{ and } i \in \mathcal{N}_D, \\ (1 - \sum_{i \in \mathcal{N}_D} q_i^e(1)) / N_U & \text{if } \sum_{i \in \mathcal{N}} q_i^e(1) < 1 \text{ and } i \in \mathcal{N}_U, \end{cases}$$

where $\hat{\theta} \in (\underline{\theta}, 1]$ is uniquely defined by $(q_1^e(\hat{\theta}), \dots, q_n^e(\hat{\theta})) \in \Delta$. Then $\mathbf{r}^* \in \Delta$ and $\mathbf{r}^*$ equalizes the agents' worst-off types under ex post efficiency at $\hat{\theta}$ if $\sum_{i \in \mathcal{N}} q_i^e(1) > 1$ and at 1 if $\sum_{i \in \mathcal{N}} q_i^e(1) \leq 1$.

Proof. By the definition of ex post efficiency, for all $i \in \mathcal{N}$ and $\theta \in [\underline{\theta}, 1]$, $q_i^e(\theta_i) = \prod_{j \in \mathcal{N} \setminus \{i\}} F_j(\theta_j)$, which is continuous and increasing in θ_i on $[\underline{\theta}, 1]$. We first show that $\sum_{i \in \mathcal{N}} q_i^e(\underline{\theta}) < 1$. Given our assumption of $n_D \geq 1$, each agent $i \in \mathcal{N}_U$ has $q_i^e(\underline{\theta}) = 0$, which implies that $\sum_{i \in \mathcal{N}} q_i^e(\underline{\theta}) = \sum_{i \in \mathcal{N}_D} q_i^e(\underline{\theta})$. So, it is sufficient to show that

$$\sum_{i \in \mathcal{N}_D} q_i^e(\underline{\theta}) < 1. \quad (\text{B.12})$$

If $n_D = 1$, then $q_i^e(\underline{\theta}) < 1$, and so (B.12) holds. If $n_D > 1$, then for $i \in \mathcal{N}_D$, $q_i^e(\underline{\theta}) = 0$, so again (B.12) holds. Thus, we have established (B.12).

Case 1: $\sum_{i \in \mathcal{N}} q_i^e(1) \geq 1$. Using (B.12), we have

$$\sum_{i \in \mathcal{N}} q_i^e(\underline{\theta}) < 1 \leq \sum_{i \in \mathcal{N}} q_i^e(1).$$

By the continuity and monotonicity of $q_i^e(\cdot)$ on $[\underline{\theta}, 1]$, there exists a unique $\hat{\theta} \in (\underline{\theta}, 1]$ such that $\sum_{i \in \mathcal{N}} q_i^e(\hat{\theta}) = 1$. Further, $q_i^e(\hat{\theta}) \in [0, 1]$ for all i . So, $(q_1^e(\hat{\theta}), \dots, q_n^e(\hat{\theta})) \in \Delta$. It then follows that $\mathbf{r}^* \in \Delta$ and, since for all $i \in \mathcal{N}$, $q_i^e(\hat{\theta}) = r_i$, all agents have a common worst-off type $\hat{\theta}$ under ex post efficiency and ownership $\mathbf{r}^*$.

Case 2: $\sum_{i \in \mathcal{N}} q_i^e(1) < 1$. In this case, we have $\sum_{j \in \mathcal{N}_D} q_j^e(1) < 1$. Also, $N_D \geq 2$ because with only one agent j in $\mathcal{N}_D$, $q_j^e(1) = 1$, which contradicts $\sum_{j \in \mathcal{N}_D} q_j^e(1) < 1$. It then follows that $\mathbf{r}^* \in \Delta$. Since for $j \in \mathcal{N}_D$, $q_j^e(1) = r_j^*$, agents in $\mathcal{N}_D$ have a worst-off type of 1. For all $i \in \mathcal{N}_U$, $q_i^e(1) = \prod_{j \in \mathcal{N}_D} F_j(1)$, which is the same for all agents in $\mathcal{N}_U$, implying that for any $i \in \mathcal{N}_U$, $N_U q_i^e(1) = \sum_{j \in \mathcal{N}_U} q_j^e(1)$. Thus, for any $i \in \mathcal{N}_U$,

$$\begin{aligned} r_i^* - q_i^e(1) &= \left(1 - \sum_{i \in \mathcal{N}_D} q_i^e(1)\right) / N_U - q_i^e(1) = \frac{1}{N_U} \left(1 - \sum_{i \in \mathcal{N}_D} q_i^e(1) - N_U q_i^e(1)\right) \\ &= \frac{1}{N_U} \left(1 - \sum_{i \in \mathcal{N}_D} q_i^e(1) - \sum_{j \in \mathcal{N}_U} q_j^e(1)\right) = \frac{1}{N_U} \left(1 - \sum_{i \in \mathcal{N}} q_i^e(1)\right) \\ &> 0. \end{aligned}$$

Because for all $\theta \in [0, 1]$, $q_i^e(1) < r_i^*$, it follows that 1 is also the worst-off type for each agent in $\mathcal{N}_U$. $\square$

Lemma B.8. *Given $\underline{\theta} \in [0, 1)$, ownership $\mathbf{r}^*$ maximizes $\Pi_c(\mathbf{r})$ subject to $\mathbf{r} \in \Delta$.*

Proof. By Lemma B.6, $\Pi_c(\mathbf{r})$ is concave, so we need only confirm that necessary conditions for a maximum are satisfied at $\mathbf{r}^*$. For each agent i , let $\hat{\theta}_i^e : [0, 1] \rightarrow [\underline{\theta}_i, \bar{\theta}_i]$ map r_i onto a worst-off type for agent i . For an agent $i \in \mathcal{N}_U$, $\hat{\theta}_i^e(r_i)$ is uniquely defined unless $r_i = 0$, in which case $\underline{\theta}$ is a worst-off type. So we can define a worst-off type function for $i \in \mathcal{N}_U$ by

$$\hat{\theta}_i^e(r_i) \equiv \begin{cases} \underline{\theta} & \text{if } r_i = 0, \\ q_i^{e^{-1}}(r_i) & \text{if } r_i \in (0, q_i^e(1)], \\ 1 & \text{if } r_i > q_i^e(1), \end{cases}$$

which is continuous and nondecreasing. In addition, it is differentiable everywhere if $\underline{\theta} = 0$ and otherwise differentiable except at $r_i = q_i^e(1) < 1$. For agent $j \in \mathcal{N}_D$, $\hat{\theta}_j^e(r_j)$ is uniquely defined unless $n_D = 1$ and $r_j = 1$, in which case 1 is a worst-off type. So we can define a

worst-off type function for $j \in \mathcal{N}_D$ by

$$\hat{\theta}_j^e(r_j) \equiv \begin{cases} \underline{\theta} & \text{if } r_j < q_j^e(\underline{\theta}), \\ q_j^{e^{-1}}(r_j) & \text{if } r_j \in [q_j^e(\underline{\theta}), 1), \\ 1 & \text{if } r_j = 1, \end{cases}$$

which is continuous, nondecreasing, and, except at $r_j = q_j^e(\underline{\theta}) < 1$, differentiable.

Note that Π_c can be written as

$$\Pi_c(\mathbf{r}) = \sum_{i \in \mathcal{N}} \left(\int_{\underline{\theta}_i}^{\hat{\theta}_i^e(r_i)} \Psi_i^S(\theta) q_i^e(\theta) dF_i(\theta) + \int_{\hat{\theta}_i^e(r_i)}^{\bar{\theta}_i} \Psi_i^B(\theta) q_i^e(\theta) dF_i(\theta) - r_i \hat{\theta}_i^e(r_i) \right).$$

Consider the maximization problem $\max_{\mathbf{r}} \Pi_c(\mathbf{r})$ subject to $\mathbf{r} \in \Delta$, i.e., subject to $r_i \geq 0$ for all $i \in \mathcal{N}$ and $\sum_{i \in \mathcal{N}} r_i = 1$. This has Lagrangian

$$\mathcal{L}(\mathbf{r}) = \Pi_c(\mathbf{r}) + \lambda \left(\sum_{i \in \mathcal{N}} r_i - 1 \right) + \sum_{i \in \mathcal{N}} \gamma_i r_i = \Pi_c(\mathbf{r}) - \lambda + \sum_{i \in \mathcal{N}} (\lambda + \gamma_i) r_i,$$

where for all $i \in \mathcal{N}$, $\gamma_i \geq 0$. Wherever $\hat{\theta}_i^e(r_i)$ is differentiable, we have

$$\frac{\partial \mathcal{L}(\mathbf{r})}{\partial r_i} = \hat{\theta}_i^{e'}(r_i) \left(q_i^e(\hat{\theta}_i^e(r_i)) - r_i \right) - \hat{\theta}_i^e(r_i) + \lambda + \gamma_i.$$

For a maximum, we require the KKT conditions that for all $i \in \mathcal{N}$, $\frac{\partial \mathcal{L}}{\partial r_i} = 0$, $\gamma_i \geq 0$, $r_i \geq 0$, and $\gamma_i r_i = 0$, and that $\sum_{i \in \mathcal{N}} r_i = 1$.

Using Lemma B.7, if $\sum_{i \in \mathcal{N}} q_i^e(1) \geq 1$, then for all $i \in \mathcal{N}$, $\hat{\theta}_i^e(r_i^*) = \hat{\theta} \in (\underline{\theta}, 1]$ and $q_i^e(\hat{\theta}) = r_i^*$, which implies that $\hat{\theta}_i^{e'}(r_i^*)$ exists and is nonnegative and that all KKT conditions are satisfied at $\mathbf{r} = \mathbf{r}^*$, $\gamma_i = 0$, and $\lambda = \hat{\theta}$. Also using Lemma B.7, if $\sum_{i \in \mathcal{N}} q_i^e(1) < 1$, then $n_D \geq 2$ and for all $i \in \mathcal{N}$, $\hat{\theta}_i^e(r_i^*) = 1$, which implies that $\hat{\theta}_i^{e'}(r_i^*)$ exists and is nonnegative. Further, for all $i \in \mathcal{N}_U$, $\hat{\theta}_i^{e'}(r_i^*) = 0$, and for all $j \in \mathcal{N}_D$, $q_j^e(1) = r_j^*$. Thus, all KKT conditions are satisfied at $\mathbf{r} = \mathbf{r}^*$, $\gamma_i = 0$, and $\lambda = 1$. This completes the proof that $\mathbf{r}^*$ maximizes $\Pi_c(\mathbf{r})$. $\square$

Lemma B.9. For $\underline{\theta} < 1$, $\Pi_c(\mathbf{r}^*) > 0$.

Proof. An efficient, IC mechanism for which the interim individual rationality constraint of each agent is binding when evaluated at the agent's worst-off type is a Vickrey-Clarke-Groves (VCG) mechanism (see Vickrey, 1961; Clarke, 1971; Groves, 1973). Standard mechanism design results imply that the VCG mechanism maximizes ex ante expected revenue subject to ex post efficiency, (dominant-strategy) incentive compatibility, and interim individual rationality.² By the payoff equivalence theorem (see, e.g., B6rgers, 2015), any ex post efficient

²The notion of incentive compatibility is not material here. The VCG mechanism being dominant-strategy incentive compatible implies that it is also Bayesian incentive compatible. With independent private values, for any ex post efficient mechanism that is Bayesian incentive compatible, there exists an equivalent mechanism that is dominant strategy incentive compatible.

mechanism that is incentive compatible and that satisfies the interim individual rationality constraints with equality for the worst-off types induces the same interim payoffs, and thus ex ante expected payoffs, as the VCG mechanism. Consequently, focusing on the VCG mechanism is without loss of generality within the context of ex post efficient mechanisms.

We denote the VCG mechanism by $\langle \mathbf{Q}^e, \mathbf{T}_r \rangle$. For a given vector of reported types $\boldsymbol{\theta}$, maximal social surplus, denoted $W(\boldsymbol{\theta})$, is

$$W(\boldsymbol{\theta}) = \max_{\hat{\mathbf{Q}} \in \Delta} \sum_{i \in \mathcal{N}} \theta_i \hat{Q}_i(\boldsymbol{\theta}) = \sum_{i \in \mathcal{N}} \theta_i Q_i^e(\boldsymbol{\theta}),$$

where the second equality holds because $\mathbf{Q}^e$ is the ex post efficient allocation rule. Given worst-off type for agent i , $\hat{\theta}_i \in [\underline{\theta}_i, \bar{\theta}_i]$, the VCG transfer from agent i is

$$T_{r,i}(\boldsymbol{\theta}) = W(\hat{\theta}_i, \boldsymbol{\theta}_{-i}) - (W(\boldsymbol{\theta}) - \theta_i Q_i^e(\boldsymbol{\theta})) - \hat{\theta}_i r_i.$$

Accordingly, the ex post budget surplus, denoted $\Pi_c(\boldsymbol{\theta}, \mathbf{r})$, is

$$\Pi_c(\boldsymbol{\theta}, \mathbf{r}) = \sum_{i \in \mathcal{N}} T_{r,i}(\boldsymbol{\theta}) = \sum_{i \in \mathcal{N}} W(\hat{\theta}_i, \boldsymbol{\theta}_{-i}) - (n-1)W(\boldsymbol{\theta}) - \sum_{i \in \mathcal{N}} \hat{\theta}_i r_i, \quad (\text{B.13})$$

where the second equality follows from substituting $T_{r,i}(\boldsymbol{\theta})$. The transfer $T_{r,i}(\boldsymbol{\theta})$ depends both directly on r_i , as is evident from the definition, and indirectly because $\hat{\theta}_i$ is also a function of r_i .

Take $\boldsymbol{\theta}$ and $\mathbf{r}$ as given. Let $\hat{\theta}_i$ be agent i 's worst-off type under the VCG mechanism. Fixing an agent $i \in \mathcal{N}$, when the type profile is $(\hat{\theta}_i, \boldsymbol{\theta}_{-i})$, social welfare under the efficient allocation for that type profile, $\mathbf{Q}^e(\hat{\theta}_i, \boldsymbol{\theta}_{-i})$, is at least as large as the social welfare under the allocation $\mathbf{Q}^e(\boldsymbol{\theta})$, because, by construction, $\mathbf{Q}^e(\hat{\theta}_i, \boldsymbol{\theta}_{-i})$ maximizes social welfare given the type profile $(\hat{\theta}_i, \boldsymbol{\theta}_{-i})$. That is, $W(\hat{\theta}_i, \boldsymbol{\theta}_{-i}) \geq W(\boldsymbol{\theta}) + (\hat{\theta}_i - \theta_i)Q_i^e(\boldsymbol{\theta})$, which is equivalent to $W(\boldsymbol{\theta}) - W(\hat{\theta}_i, \boldsymbol{\theta}_{-i}) \leq (\theta_i - \hat{\theta}_i)Q_i^e(\boldsymbol{\theta})$. Summing up yields

$$\sum_{i \in \mathcal{N}} \left(W(\boldsymbol{\theta}) - W(\hat{\theta}_i, \boldsymbol{\theta}_{-i}) \right) \leq W(\boldsymbol{\theta}) - \sum_{i \in \mathcal{N}} \hat{\theta}_i Q_i^e(\boldsymbol{\theta}). \quad (\text{B.14})$$

Using (B.13) and (B.14), we have

$$\begin{aligned} \Pi_c(\boldsymbol{\theta}, \mathbf{r}) &= \sum_{i \in \mathcal{N}} W(\hat{\theta}_i, \boldsymbol{\theta}_{-i}) - (n-1)W(\boldsymbol{\theta}) - \sum_{i \in \mathcal{N}} \hat{\theta}_i r_i \\ &= W(\boldsymbol{\theta}) - \sum_{i \in \mathcal{N}} \hat{\theta}_i Q_i^e(\boldsymbol{\theta}) - \sum_{i \in \mathcal{N}} \left(W(\boldsymbol{\theta}) - W(\hat{\theta}_i, \boldsymbol{\theta}_{-i}) \right) \\ &\quad + \sum_{i \in \mathcal{N}} \hat{\theta}_i Q_i^e(\boldsymbol{\theta}) - \sum_{i \in \mathcal{N}} \hat{\theta}_i r_i \\ &\geq \sum_{i \in \mathcal{N}} \hat{\theta}_i Q_i^e(\boldsymbol{\theta}) - \sum_{i \in \mathcal{N}} \hat{\theta}_i r_i. \end{aligned} \quad (\text{B.15})$$

Let $\hat{\theta}_i^*$ be agent i 's worst-off type under the VCG mechanism $\langle \mathbf{Q}^e, \mathbf{T}_{r^*} \rangle$, where $\mathbf{r}^*$ is as

defined in Lemma B.7. We now show that $\mathbb{E}_{\boldsymbol{\theta}} \left[\sum_{i \in \mathcal{N}} \hat{\theta}_i^* Q_i^e(\boldsymbol{\theta}) - \sum_{i \in \mathcal{N}} \hat{\theta}_i^* r_i^* \right] > 0$, which, using (B.15), implies that $\Pi_c(\mathbf{r}^*) = \mathbb{E}_{\boldsymbol{\theta}} [\Pi_c(\boldsymbol{\theta}, \mathbf{r}^*)] > 0$.

Using Lemma B.7, agents have a common worst-off type under $\mathbf{r}^*$ and $\mathbf{Q}^e$, i.e., $\hat{\theta}_i^* = \hat{\theta}$ for all $i \in \mathcal{N}$. Thus,

$$\Pi_c(\boldsymbol{\theta}, \mathbf{r}^*) \geq \sum_{i \in \mathcal{N}} \hat{\theta}_i^* Q_i^e(\boldsymbol{\theta}) - \sum_{i \in \mathcal{N}} \hat{\theta}_i^* r_i^* = \hat{\theta} \sum_{i \in \mathcal{N}} (Q_i^e(\boldsymbol{\theta}) - r_i^*) = 0 \quad (\text{B.16})$$

because $\sum_{i \in \mathcal{N}} Q_i(\boldsymbol{\theta}) = 1$. The inequality in (B.16) is strict for a positive measure of types (to see this, note that (B.16) holds with equality only if $Q_i^e(\boldsymbol{\theta}) = Q_i^e(\hat{\theta}, \boldsymbol{\theta}_{-i})$, which only holds for a restricted set of $\boldsymbol{\theta}$ given that $\underline{\theta} < 1$ and $f_i > 0$ for all i on $(\underline{\theta}, 1)$), and so $\Pi_c(\mathbf{r}^*) > 0$ follows. $\square$

Combining all the elements above completes the proof of Lemma A.2. $\blacksquare$

B.6 Proof of Proposition 3(iii)

Proof of Proposition 3(iii). We assume that Ψ_i^S and Ψ_i^B are increasing, and so invertible. We begin by focusing on the case with $w = 1$ and proving that $U_{1r}(r, 1) > 0$ for r sufficiently close to 0, and that $U_{1r}(r, 1) < 0$ for r sufficiently close to 1, when $F_1 = F_2$. For $w = 1$, the interim expected allocations are:

$$q_1(\theta_1) = \begin{cases} F_2(\Psi_2^{S^{-1}}(\theta_1)) & \text{if } 0 \leq \theta_1 \leq F_1^{-1}(1 - r), \\ F_2(\Psi_2^{B^{-1}}(\theta_1)) & \text{if } F_1^{-1}(1 - r) < \theta_1 \leq 1, \end{cases}$$

and

$$q_2(\theta_2) = \begin{cases} F_1(\Psi_2^S(\theta_2)) & \text{if } 0 \leq \theta_2 < \Psi_2^{S^{-1}}(F_1^{-1}(1 - r)), \\ 1 - r & \text{if } \Psi_2^{S^{-1}}(F_1^{-1}(1 - r)) \leq \theta_2 \leq \Psi_2^{B^{-1}}(F_1^{-1}(1 - r)), \\ F_1(\Psi_2^B(\theta_2)) & \text{if } \Psi_2^{B^{-1}}(F_1^{-1}(1 - r)) < \theta_2 \leq 1. \end{cases}$$

Using the characterization of worst-off types from Cramton et al. (1987), this gives us a worst-off type for agent 1 of

$$\omega_1 = \begin{cases} \Psi_2^B(F_2^{-1}(r)) & \text{if } F_2(\Psi_2^{B^{-1}}(F_1^{-1}(1 - r))) \leq r, \\ F_1^{-1}(1 - r) & \text{if } F_2(\Psi_2^{S^{-1}}(F_1^{-1}(1 - r))) < r < F_2(\Psi_2^{B^{-1}}(F_1^{-1}(1 - r))), \\ \Psi_2^S(F_2^{-1}(r)) & \text{if } r \leq F_2(\Psi_2^{S^{-1}}(F_1^{-1}(1 - r))), \end{cases}$$

and for agent 2, any type in $[\Psi_2^{S^{-1}}(F_1^{-1}(1 - r)), \Psi_2^{B^{-1}}(F_1^{-1}(1 - r))]$ is worst-off, including $\omega_2 = F_1^{-1}(1 - r)$.

We have $u_i(\theta) = \theta(q_i(\theta) - r_i) - m_i(\theta)$ and, in the bargaining mechanism,

$$m_i(\theta) = \theta(q_i(\theta) - r_i) - \int_{\omega_i}^{\theta_i} (q_i(y) - r_i)dy - \eta_i \pi(\mathbf{Q}, \boldsymbol{\omega}),$$

where $\pi(\mathbf{Q}, \boldsymbol{\omega}) = \sum_{i \in \mathcal{N}} \mathbb{E}[\Psi_i(\theta_i, \omega_i) q_i(\theta_i)] - \sum_{i \in \mathcal{N}} \omega_i r_i$.

For the case of $n = 2$, we have $(\eta_1, \eta_2) = (1, 0)$ if $w > 1/2$, $(\eta_1, \eta_2) = (0, 1)$ if $w < 1/2$, and $(\eta_1, \eta_2) = (\eta, 1 - \eta)$ for some $\eta \in [0, 1]$ if $w = 1/2$. Thus, for $w = 1$, we have $\eta_1 = 1$ and so

$$u_1(\theta_1) = \int_{\omega_1}^{\theta_1} (q_1(y) - r)dy + \sum_{i=1}^2 \mathbb{E}_{\theta_i}[\Psi_i(\theta_i, \omega_i) q_i(\theta_i)] - \omega_1 r - \omega_2(1 - r).$$

Thus, agent 1's expected net payoff when $w = 1$ is

$$U_1(r, 1) = \int_0^1 \int_{\omega_1}^{\theta_1} (q_1(y) - r)dy dF_1(\theta_1) + \sum_{i=1}^2 \int_0^1 \Psi_i(\theta_i, \omega_i) q_i(\theta_i) dF_i(\theta_i) - \omega_1 r - \omega_2(1 - r).$$

We can write the double integral in the expression for $U_1(r, 1)$ as

$$\begin{aligned} & \int_{\omega_1}^1 \int_{\omega_1}^{\theta_1} (q_1(y) - r)dy dF_1(\theta_1) - \int_0^{\omega_1} \int_{\theta_1}^{\omega_1} (q_1(y) - r)dy dF_1(\theta_1) \\ &= \int_{\omega_1}^1 (1 - F_1(y))(q_1(y) - r)dy - \int_0^{\omega_1} F_1(y)(q_1(y) - r)dy. \end{aligned}$$

Taking the case of r sufficiently close to 1 such that $\omega_1 = \Psi_2^B(F_2^{-1}(r)) > F_1^{-1}(1 - r)$, we can rewrite this as

$$\int_{\omega_1}^1 (1 - F_1(y))(F_2(\Psi_2^{B^{-1}}(y)) - r)dy - \int_0^{F_1^{-1}(1-r)} F_1(y)(F_2(\Psi_2^{S^{-1}}(y)) - r)dy - \int_{F_1^{-1}(1-r)}^{\omega_1} F_1(y)(F_2(\Psi_2^{B^{-1}}(y)) - r)dy,$$

whose derivative with respect to r , evaluated at $r = 1$, is $1 - \mathbb{E}[\theta_1]$. (Generally, for r sufficiently close to 1, we have $\omega_1 - \mathbb{E}[\theta_1] - \omega_1'(1 - F_1(\omega_1))(F_2(\Psi_2^{B^{-1}}(\omega_1)) - r) + F_1^{-1'}(1 - r)(1 - r)[F_2(\Psi_2^{S^{-1}}(F_1^{-1}(1 - r))) - F_2(\Psi_2^{B^{-1}}(F_1^{-1}(1 - r)))]$.) Turning to the summation term in the expression for $U_1(r, 1)$, for agent 1, and for r sufficiently close to 1 such that $\omega_1 = \Psi_2^B(F_2^{-1}(r)) > F_1^{-1}(1 - r)$, this is

$$\begin{aligned} & \int_0^{\omega_1} \Psi_i^S(\theta_i) q_i(\theta_i) dF_i(\theta_i) + \int_{\omega_i}^1 \Psi_i^B(\theta_i) q_i(\theta_i) dF_i(\theta_i) \\ &= \int_0^{F_1^{-1}(1-r)} \Psi_1^S(\theta_1) F_2(\Psi_2^{S^{-1}}(\theta_1)) dF_1(\theta_1) + \int_{F_1^{-1}(1-r)}^{\omega_1} \Psi_1^S(\theta_1) F_2(\Psi_2^{B^{-1}}(\theta_1)) dF_1(\theta_1) \\ & \quad + \int_{\omega_1}^1 \Psi_1^B(\theta_1) F_2(\Psi_2^{B^{-1}}(\theta_1)) dF_1(\theta_1), \end{aligned}$$

whose derivative with respect to r , evaluated at $r = 1$, is $2/f_2(1)$. (Generally, for r sufficiently close to 1, we have $\Psi_1^S(F_1^{-1}(1 - r))[F_2(\Psi_2^{B^{-1}}(F_1^{-1}(1 - r))) - F_2(\Psi_2^{S^{-1}}(F_1^{-1}(1 - r)))] + \Psi_2^{B'}(F_2^{-1}(r))F_2^{-1'}(r)r$.) For agent 2, noting that we are working with r such that

$\Psi_2^{B^{-1}}(F_1^{-1}(1-r)) \leq F_2^{-1}(r)$, the summation term for agent 2 is

$$\int_0^{\Psi_2^{S^{-1}}(F_1^{-1}(1-r))} \Psi_2^S(\theta_2)(F_1(\Psi_2^S(\theta_2)) - (1-r))dF_2(\theta_2) + \int_{\Psi_2^{B^{-1}}(F_1^{-1}(1-r))}^1 \Psi_2^B(\theta_2)(F_1(\Psi_2^B(\theta_2)) - (1-r))dF_2(\theta_2) + (1-r)\omega_2.$$

Differentiating with respect to r and evaluating at $r = 1$, we get $\int_{\Psi_2^{B^{-1}}(0)}^1 \Psi_2^B(\theta_2)dF_2(\theta_2)$. (For r sufficiently close to 1, we get $\int_0^{\Psi_2^{S^{-1}}(F_1^{-1}(1-r))} \Psi_2^S(\theta_2)dF_2(\theta_2) + \int_{\Psi_2^{B^{-1}}(F_1^{-1}(1-r))}^1 \Psi_2^B(\theta_2)dF_2(\theta_2) - \omega_2 + (1-r)\omega_2'$.) Finally, note that for r sufficiently close to 1, the derivative of $-\omega_1 r - \omega_2(1-r)$, evaluated at $r = 1$, is equal to $-1 - \frac{2}{f_2(1)}$. Thus, gathering the terms calculated above, for r sufficiently close to 1, the derivative of $U_1(r, 1)$ taken with respect to r and evaluated at $r = 1$, is

$$\begin{aligned} U_{1r}(1, 1) &= 1 - \mathbb{E}[\theta_1] + 2/f_2(1) + \int_{\Psi_2^{B^{-1}}(0)}^1 \Psi_2^B(\theta_2)dF_2(\theta_2) - 1 - 2/f_2(1) \\ &= -\int_0^1 \theta_1 dF_1(\theta_1) + \int_{\Psi_2^{B^{-1}}(0)}^1 \theta_2 dF_2(\theta_2) - \int_{\Psi_2^{B^{-1}}(0)}^1 (1 - F_2(\theta_2)) d\theta_2. \end{aligned}$$

If $F_1 = F_2 = F$, then we have

$$U_{1r}(1, 1) = -\int_0^{\Psi^{B^{-1}}(0)} \theta_1 dF(\theta_1) - \int_{\Psi^{B^{-1}}(0)}^1 (1 - F(\theta_2)) d\theta_2 < 0.$$

By analogous calculations, $U_{1r}(0, 1) > 0$ if $F_1 = F_2$. By continuity, $U_{1r}(r, 1) < 0$ for r sufficiently close to 1 and $U_{1r}(r, 1) > 0$ for r sufficiently close to 0, assuming that $F_1 = F_2$.

For $w = 0$, we can reverse the roles of agents 1 and 2 in the analysis above and replace r with $1-r$, giving us the result that $U_{2r}(r, 0) < 0$ for all r sufficiently large and $U_{2r}(r, 0) > 0$ for all r sufficiently small, assuming that $F_1 = F_2$.

Now turn to effects for the agent without the bargaining power, starting with $w = 1$. For agent 2, we have $\eta_2 = 0$ and so $u_2(\theta_2) = \int_{\omega_2}^{\theta_2} (q_2(y) - (1-r))dy$ and $U_2(r, 1) = \int_0^1 \int_{\omega_2}^{\theta_2} (q_2(y) - (1-r))dy dF_2(\theta_2)$. This is analogous to the double integral term analyzed above. Replacing 1 with 2 and r with $1-r$ in the expression above, we have $U_2(r, 1) = \int_{\omega_2}^1 (1 - F_2(y))(q_2(y) - (1-r))dy - \int_0^{\omega_2} F_2(y)(q_2(y) - (1-r))dy$. Thus, using the definition of q_2 , we have

$$U_2(r, 1) = \int_{\Psi_2^{B^{-1}}(F_1^{-1}(1-r))}^1 (1 - F_2(y))(F_1(\Psi_2^B(y)) - (1-r))dy - \int_0^{\Psi_2^{S^{-1}}(F_1^{-1}(1-r))} F_2(y)(F_1(\Psi_2^S(y)) - (1-r))dy.$$

Differentiating with respect to r , we get

$$\begin{aligned} U_{2r}(r, 1) &= \int_{\Psi_2^{B^{-1}}(F_1^{-1}(1-r))}^1 (1 - F_2(y))dy - \int_0^{\Psi_2^{S^{-1}}(F_1^{-1}(1-r))} F_2(y)dy \\ &= \begin{cases} -\int_0^{\Psi_2^{S^{-1}}(1)} F_2(y)dy < 0 & \text{if } r = 0, \\ \int_{\Psi_2^{B^{-1}}(0)}^1 (1 - F_2(y))dy > 0 & \text{if } r = 1. \end{cases} \end{aligned}$$

By continuity, $U_{2r}(r, 1) > 0$ for all r sufficiently close to 1 and $U_{2r}(r, 1) < 0$ for all r sufficiently close to 0. Analogously, $U_{1r}(r, 0) > 0$ for r sufficiently close to 1 and $U_{1r}(r, 0) < 0$ for r sufficiently close to 0.

for r sufficiently close to 0. ■

B.7 Proof of Proposition 7

Proof of Proposition 7. For every pair i and j of agents in $\mathcal{N}_D$, it is possible for either to have the highest type, so the ranking of those agents matters for efficiency. In addition, for $\underline{\theta} < 1$, for every pair of agents, one in $\mathcal{N}_U$ and one in $\mathcal{N}_D$, it is possible for either to have the highest type, so the ranking of those agents matters for efficiency.

We complete the proof with two lemmas. In Lemma B.10, we show that for two agents with overlapping supports, the ranking of their types cannot be the same as the ranking of their ironed weighted virtual types for all types in an open interval of the support overlap if the bargaining weights differ. This implies that for ex post efficiency, all agents in $\mathcal{N}_D$ must have the same bargaining weight and that when $\underline{\theta} < 1$, all agents must have the same bargaining weight. Then we address the case of $\underline{\theta} \geq 1$ in Lemma B.11. It is convenient as part of Lemma B.11 to also state an additional result for the case of one agent in $\mathcal{N}_D$, and so the proof of this lemma also serves as the proof of Proposition 2, which assumes $n_U = 1$ and $n_D = 1$.

Lemma B.10. *The ranking of the ironed weighted virtual type functions for two agents cannot be the same as the ranking of their types for all types in an open interval in the support of both agents' type distributions if the agents' bargaining weights differ.*

Proof. Suppose that $w_i \neq w_j$ and let A be an open interval that is a subset of the interval of overlap in the supports of the distributions of agents i and j . Suppose that for all $\theta \in A$, the ranking of agents i and j according to their types is the same as the ranking of those agents according to their weighted ironed virtual types:

$$\bar{\Psi}_{j, \frac{w_j}{\rho}}(\theta_j, \omega_j) > \bar{\Psi}_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i) \Leftrightarrow \theta_j > \theta_i.$$

If there exists $\theta \in A$ such that $\bar{\Psi}_{j, w_j/\rho}(\theta, \omega_j) > \bar{\Psi}_{i, w_i/\rho}(\theta, \omega_i)$, then by continuity, there exists $\varepsilon > 0$ sufficiently small that $\bar{\Psi}_{j, w_j/\rho}(\theta, \omega_j) > \bar{\Psi}_{i, w_i/\rho}(\theta + \varepsilon, \omega_i)$, violating the requirement of a common ranking because $\theta < \theta + \varepsilon$. A similar contradiction arises if there exists θ such that $\bar{\Psi}_{j, w_j/\rho}(\theta, \omega_j) < \bar{\Psi}_{i, w_i/\rho}(\theta, \omega_i)$. Thus, we require that for all $\theta \in A$, $\bar{\Psi}_{j, w_j/\rho}(\theta, \omega_j) = \bar{\Psi}_{i, w_i/\rho}(\theta, \omega_i)$.

If the ironed weighted virtual type functions are constant and equal on any open subset of A , then, again, the ranking is not preserved. So it must be that they are increasing on A . Given that the ironed weighted virtual type functions are increasing on A , they are either equal to an unironed portion of the weighted virtual value or an unironed portion of the

weighted virtual cost, and, as we now show, these cannot be the same on an open interval if the weights differ.

To see this, suppose that for $\theta \in A$, $\bar{\Psi}_{j,w_j/\rho}(\theta, \omega_j) = \Psi_{j,w_j/\rho}^B(\theta) = \bar{\Psi}_{i,w_i/\rho}(\theta, \omega_i) = \Psi_{i,w_i/\rho}^S(\theta)$. Then we have

$$\theta - (1 - \frac{w_j}{\rho}) \frac{1 - F_j(\theta)}{f_j(\theta)} = \theta + (1 - \frac{w_i}{\rho}) \frac{F_i(\theta)}{f_i(\theta)},$$

which we can rewrite as

$$\frac{F_i(\theta) f_j(\theta)}{(1 - F_j(\theta)) f_i(\theta) + F_i(\theta) f_j(\theta)} = \frac{\rho - w_j}{w_i - w_j} \equiv C$$

and, alternatively, as

$$\frac{(1 - F_j(\theta)) f_i(\theta)}{F_i(\theta) f_j(\theta) + (1 - F_j(\theta)) f_i(\theta)} = \frac{\rho - w_i}{w_j - w_i} \equiv C'$$

Using $\rho \geq \max\{w_j, w_i\}$, if $w_j < w_i$, then $C \geq 1$, and we require that for all $\theta \in A$, $\frac{1 - F_j(\theta)}{f_j(\theta)} / \frac{F_i(\theta)}{f_i(\theta)} = \frac{1 - C}{C}$, which is a contradiction because the left side is positive and $(1 - C)/C \leq 0$; and if $w_j > w_i$, then $C' \geq 1$, and we require that $\frac{F_i(\theta)}{f_i(\theta)} / \frac{(1 - F_j(\theta))}{f_j(\theta)} = \frac{1 - C'}{C'}$, which is similarly a contradiction. An analogous contradiction arises if we start from the supposition that for all $\theta \in A$, $\bar{\Psi}_{j,w_j/\rho}(\theta, \omega_j) = \Psi_{j,w_j/\rho}^S(\theta) = \bar{\Psi}_{i,w_i/\rho}(\theta, \omega_i^*) = \Psi_{i,w_i/\rho}^B(\theta)$.

Now suppose that

$$\bar{\Psi}_{j,w_j/\rho}(\theta, \omega_j) = \Psi_{j,w_j/\rho}^B(\theta) = \bar{\Psi}_{i,w_i/\rho}(\theta, \omega_i) = \Psi_{i,w_i/\rho}^B(\theta).$$

Then we require that

$$\theta - (1 - \frac{w_j}{\rho}) \frac{1 - F_j(\theta)}{f_j(\theta)} = \theta - (1 - \frac{w_j}{\rho}) \frac{1 - F_i(\theta)}{f_i(\theta)}.$$

We can write this as

$$\frac{1 - F_i(\theta)}{f_i(\theta)} / \left(\frac{1 - F_j(\theta)}{f_j(\theta)} - \frac{1 - F_i(\theta)}{f_i(\theta)} \right) = \frac{\rho - w_j}{w_j - w_i} \equiv C'',$$

which implies that $\frac{1 - F_j(\theta)}{f_j(\theta)} / \frac{1 - F_i(\theta)}{f_i(\theta)} = \frac{1 + C''}{C''}$, giving us a contradiction for $w_j < w_i$ because the left side is positive and the right side is negative. An analogous contradiction obtains if $w_j > w_i$. The remaining case, $\bar{\Psi}_{j,w_j/\rho}(\theta, \omega_j) = \Psi_{j,w_j/\rho}^S(\theta) = \bar{\Psi}_{i,w_i/\rho}(\theta, \omega_i) = \Psi_{i,w_i/\rho}^S(\theta)$, provides an analogous contradiction, which completes the proof. $\square$

Lemma B.11. *Let $\bar{w}_U \equiv \max_{\ell \in \mathcal{N}_U \text{ s.t. } r_\ell > 0} w_\ell$, let agent $i \in \mathcal{N}_U$ be such that $r_i > 0$, and assume constant marginal values and $\underline{\theta} \geq 1$. Ex post efficiency requires that there exist $j \in \mathcal{N}_D$ such that $\bar{\Psi}_{i, \frac{w_i}{\bar{w}}}^S(1) \leq \bar{\Psi}_{j, \frac{w_D}{\bar{w}}}(\underline{\theta}, \omega_j)$, which is satisfied if $w_i = \bar{w}_U = w_D$. If, in addition, virtual cost and virtual value functions are monotone, then ex post efficiency requires that $\frac{\max\{\bar{w}_U, w_D\} - w_i}{\max\{\bar{w}_U, w_D\}} \leq (\underline{\theta} - 1) f_i(1)$, and if there is only one agent j in $\mathcal{N}_D$, then ex post efficiency*

is possible if and only if $1 + (1 - \frac{w_i}{\max\{w_i, w_j\}})/f_i(1) \leq \underline{\theta} - (1 - \frac{w_j}{\max\{w_i, w_j\}})/f_j(\underline{\theta})$.

Proof. Given $\underline{\theta} \geq 1$, ex post efficiency requires that resources move from agent $i \in \mathcal{N}_U$ with $r_i > 0$ to an agent in $\mathcal{N}_D$ for all type realizations. Thus, we require that for some agent $j \in \mathcal{N}_D$,

$$\bar{\Psi}_{i, \frac{w_i}{\rho}}(\theta_i, \omega_i) \leq \bar{\Psi}_{j, \frac{w_j}{\rho}}(\theta_j, \omega_j), \quad (\text{B.17})$$

where ω_ℓ is agent ℓ 's worst-off type and ρ is the Lagrange multiplier on the no-deficit constraint. Under ex post efficiency, agent i is a seller with worst-off type 1, all agents in $\mathcal{N}_D$ have common bargaining weight w_D (as just shown), and $\rho = \bar{w} \equiv \max\{\bar{w}_U, w_D\}$. So, (B.17) becomes $\bar{\Psi}_{i, \frac{w_i}{\bar{w}}}^S(\theta_i) \leq \bar{\Psi}_{j, \frac{w_D}{\bar{w}}}(\theta_j, \omega_j)$, and this must hold even for $\theta_i = 1$ and even if all agents in $\mathcal{N}_D$ have type $\underline{\theta}$. Thus, we require that there exists $j \in \mathcal{N}_D$ such that

$$\bar{\Psi}_{i, \frac{w_i}{\bar{w}}}^S(1) \leq \bar{\Psi}_{j, \frac{w_D}{\bar{w}}}(\underline{\theta}, \omega_j). \quad (\text{B.18})$$

This is satisfied if $w_i = \bar{w}_U = w_D$ because then (B.18) becomes $1 \leq \underline{\theta}$, which holds. (It need not hold for, say, $w_i = \bar{w}_U = 1$ and $w_D = 0$ because then (B.18) becomes $1 \leq \bar{\Psi}_{j,0}(\underline{\theta}, \omega_j)$, which if, say, $\omega_j = \underline{\theta}$ amounts to $1 \leq \bar{\Psi}_j^B(\underline{\theta})$, which can fail to hold (e.g., for any distribution and $\underline{\theta}$ sufficiently close to 1).

Assuming monotone virtual cost and virtual value functions, (B.18) implies that

$$1 + \frac{1 - \frac{w_i}{\bar{w}}}{f_i(1)} = \Psi_{i, \frac{w_i}{\bar{w}}}^S(1) \leq \bar{\Psi}_{j, \frac{w_D}{\bar{w}}}(\underline{\theta}, \omega_j) \leq \underline{\theta}, \quad (\text{B.19})$$

where the final inequality uses $\bar{\Psi}_{j, \frac{w_D}{\bar{w}}}(\underline{\theta}, \omega_j) \in [\Psi_{j, \frac{w_D}{\bar{w}}}^B(\underline{\theta}), \Psi_{j, \frac{w_D}{\bar{w}}}^S(\underline{\theta})] = [\Psi_{j, \frac{w_D}{\bar{w}}}^B(\underline{\theta}), \underline{\theta}]$. Thus, for all agents $i \in \mathcal{N}_U$ with $r_i > 0$, $\frac{\bar{w} - w_i}{\bar{w}} \leq (\underline{\theta} - 1) f_i(1)$, as stated in the lemma.

Continuing with the assumption of monotone virtual cost and virtual value functions, if there is only one agent j in $\mathcal{N}_D$, then under ex post efficiency, agent j is a buyer with worst-off type $\underline{\theta}$, and so $\bar{\Psi}_{j, w_D/\bar{w}}(\underline{\theta}, \omega_j) = \Psi_{j, w_D/\bar{w}}^B(\underline{\theta}) = \underline{\theta} - \frac{1 - \frac{w_D}{\bar{w}}}{f_j(\underline{\theta})}$, which gives us

$$1 + \frac{1 - \frac{w_i}{\bar{w}}}{f_i(1)} \leq \underline{\theta} - \frac{1 - \frac{w_D}{\bar{w}}}{f_j(\underline{\theta})}. \quad (\text{B.20})$$

Under condition (B.20), trade occurs for all type realizations, and so the expected budget surplus under ex post efficiency and binding IR for the agents' worst-off types is

$$\begin{aligned} \Pi_c(\mathbf{r}) &\equiv \mathbb{E}[\Psi_j^B(\theta_j)q_j^e(\theta_j) + \sum_{i \in \mathcal{N} \setminus \{j\}} \Psi_i^S(\theta_i)q_i^e(\theta_i)] - r_j \underline{\theta} - \sum_{i \in \mathcal{N} \setminus \{j\}} r_i \\ &= \underline{\theta} - r_j \underline{\theta} - \sum_{i \in \mathcal{N} \setminus \{j\}} r_i \\ &\geq 0, \end{aligned}$$

where the second equality uses $q_j^e(\theta_j) = 1$, $q_i^e(\theta_i) = 0$ for $i \neq j$, and $\mathbb{E}[\Psi_j^B(\theta_j)] = \underline{\theta}$, and the inequality uses $\underline{\theta} \geq 1$ and $\sum_{i \in \mathcal{N}} r_i = 1$. Because Π_c is nonnegative, it follows that (B.20) is necessary and sufficient for ex post efficiency. Efficient trade can be achieved in an IC, IR,

no-deficit way and that allocates Π_c to the agent with the higher bargaining weight using a posted price of 1 if $w_i > w_j$ and $\underline{\theta}$ if $w_j > w_i$; and if $w_i = w_j$, then Π_c can be allocated with share $\eta_\ell \in [0, 1]$ going to agent ℓ using the appropriate posted price in $[1, \underline{\theta}]$. $\square$

Combining these results completes the proof of Proposition 7. To provide further characterization of the set of bargaining weights consistent with ex post efficiency, we conclude with two additional lemmas. Recall from Theorem 2 that for $\underline{\theta} \geq 1$, $\mathcal{W}_c^0(\underline{\theta})$ is the interval $[\underline{w}(\underline{\theta}), \bar{w}(\underline{\theta})]$. For $n = 2$, $\underline{\theta} \geq 1$, and monotone virtual cost and virtual value functions, we can write equation (7) for bargaining weights $(w, 1 - w)$ with $w \geq 1/2$ and again for $w < 1/2$. These two inequalities imply that

$$\underline{w}(\underline{\theta}) = \max\left\{0, \frac{1 - (\underline{\theta} - 1)f_1(1)}{2 - (\underline{\theta} - 1)f_1(1)}\right\}.$$

and

$$\bar{w}(\underline{\theta}) = \min\left\{1, \frac{1}{2 - (\underline{\theta} - 1)f_2^P(0)}\right\}.$$

Using these, we have:

Lemma B.12. *For $n = 2$, $\underline{\theta} > 1$, and monotone virtual cost and virtual value functions, $\underline{w}(\underline{\theta})$ is decreasing and strictly concave when it is greater than 0.*

Proof. Assume that $\underline{\theta} \geq 1$. We focus on $\underline{\theta}$ such that $\underline{w}(\underline{\theta}) > 0$, which implies that $1 - (\underline{\theta} - 1)f_1(1) > 0$. Differentiating $\underline{w}(\underline{\theta})$, we have:

$$\underline{w}'(\underline{\theta}) = \frac{-f_1(1)}{(2 - (\underline{\theta} - 1)f_1(1))^2} < 0.$$

Differentiating again,

$$\underline{w}''(\underline{\theta}) = \frac{-2(f_1(1))^2}{(2 - (\underline{\theta} - 1)f_1(1))^3} < 0.$$

Thus, $\underline{w}(\underline{\theta})$ is decreasing and strictly concave. $\blacksquare$

Lemma B.13. *For $n = 2$, $\underline{\theta} > 1$, and monotone virtual cost and virtual value functions, $\bar{w}(\underline{\theta})$ is increasing and strictly convex when it is less than 1.*

Proof. Assume that $\underline{\theta} > 1$. We focus on $\underline{\theta}$ such that $\bar{w}(\underline{\theta}) < 1$, which implies that $f_2^P(0) < \frac{1}{\underline{\theta}-1}$. Differentiating $\bar{w}(\underline{\theta})$, we have

$$\bar{w}'(\underline{\theta}) = \frac{f_2^P(0)}{(2 - (\underline{\theta} - 1)f_2^P(0))^2} > 0.$$

Differentiating again,

$$\bar{w}''(\underline{\theta}) = \frac{2(f_2^P(0))^2}{(2 - (\underline{\theta} - 1)f_2^P(0))^3} > 0.$$

Thus, $\bar{w}(\underline{\theta})$ is increasing and strictly convex. $\square$

This concludes the proof of Proposition 7. $\blacksquare$

B.8 Proof of Theorem 3

Proof of Theorem 3. For decreasing marginal values, given $\underline{\theta} \geq 1$ and $r \in [0, 1]$, let $\Pi_d(r)$ denote the expected budget surplus under ex post efficiency and binding IR for agents' worst-off types. It then follows that ex post efficiency is possible with $w = 1/2$ if and only if $\Pi_d(r) \geq 0$, where Π_d can be written as

$$\Pi_d(r) = \sum_{i=1}^2 \mathbb{E}_{\theta} \left[\Psi_i(\theta_i, \omega_i^e(r)) Q_i^e(\theta) - \frac{Q_i^e(\theta)^2}{2} \right] - \omega_1^e(r)r + r^2/2 - \omega_2^e(r)(1-r) + (1-r)^2/2.$$

where $\omega_i^e(r)$ is a worst-off type for agent i under ex post efficiency.

We proceed with a series of lemmas. The first establishes the “shape” of Π_d .

Lemma B.14. (i) If $\underline{\theta} = 1$, then $\Pi_d(r)$ is quasiconcave in r with a unique interior maximum at $r = \frac{1+\mu_1^P-\mu_2^P}{2}$. (ii) If $\underline{\theta} \in (1, 2)$, then $\Pi_d(r)$ is increasing at $r = 1$, has a unique interior maximum at $r^* = \frac{1+\omega_1^*- \omega_2^*}{2}$, where $q_1^e(\omega_1^*) = r^* = 1 - q_2^e(\omega_2^*)$ and where r^* is decreasing in $\underline{\theta}$, has a unique interior minimum at $r = \frac{3-\underline{\theta}}{2} \in (r^*, 1)$, and is increasing at $r = 1$.

Proof. First, we confirm that there exist interior worst-off types $\omega_1^* \in (1, 2)$ and $\omega_2^* \in (\underline{\theta}, 1+\underline{\theta})$ such that $q_1^e(\omega_1^*) = \frac{1+\omega_1^*-\omega_2^*}{2} = 1 - q_2^e(\omega_2^*)$. Because $q_1^e(\theta_1) = \int_1^2 \max\{0, \frac{1+\theta_1-(x+\underline{\theta}-1)}{2}\} dF_2^P(x)$ and $q_2^e(\theta_2) = \int_1^2 \min\{1, \frac{1-\theta_1+\theta_2}{2}\} dF_1(\theta_1)$, we have

$$\frac{1 + \omega_1^* - \omega_2^*}{2} = q_1^e(\omega_1^*) = \int_0^{\min\{1, 1+\omega_1^*-\underline{\theta}\}} \frac{1 + \omega_1^* - x - \underline{\theta}}{2} dF_2^P(x)$$

and

$$\frac{1 - \omega_1^* + \omega_2^*}{2} = q_2^e(\omega_2^*) = \int_{\max\{1, \omega_2^*-1\}}^2 \frac{1 - x + \omega_2^*}{2} dF_1(x) + F_1(\omega_2^* - 1),$$

where, of course, $F_1(\omega_2^* - 1) = 0$ if $\omega_2^* \leq 2$ and $f_1(x) = 0$ for $x < 1$. Solving these two equations, we have

$$(\omega_1^*, \omega_2^*) = \begin{cases} (\mu_1, \mu_2^P + \underline{\theta}) & \text{if } \underline{\theta} \leq \min\{\mu_1, 2 - \mu_2^P\}, \\ \left(2 - \int_{\mu_2^P+\underline{\theta}-1}^2 F_1(x)dx, \mu_2^P + \underline{\theta}\right) & \text{if } 2 - \mu_2^P < \underline{\theta} \leq \mu_1, \\ \left(\mu_1, 1 + \mu_1 - \int_0^{1+\mu_1-\underline{\theta}} F_2^P(x)dx\right) & \text{if } \mu_1 < \underline{\theta} \leq 2 - \mu_2^P, \\ \left(2 - \int_{\omega_2^*-1}^2 F_1(x)dx, 1 + \omega_1^* - \int_0^{1+\omega_1^*-\underline{\theta}} F_2^P(x)dx\right) & \text{if } \max\{\mu_1, 2 - \mu_2^P\} < \underline{\theta}, \end{cases}$$

where the final row offers only a system of equations and not an explicit solution. Nevertheless, it is straightforward to confirm that in all cases, $\omega_1^* \in (1, 2)$ and $\omega_2^* \in (\underline{\theta}, 1 + \underline{\theta})$. Given

this and $r^* = q_1^e(\omega_1^*)$, it follows that $r^* \in (0, 1)$ and that $r^* < \frac{3-\theta}{2}$. Further, we have

$$\frac{\partial \omega_1^*}{\partial \underline{\theta}} = \begin{cases} 0 & \text{if } \underline{\theta} \leq 2 - \mu_2^P, \\ \frac{\partial \omega_2^*}{\partial \underline{\theta}} F_1(\omega_2^* - 1) & \text{if } \underline{\theta} > 2 - \mu_2^P, \end{cases}$$

and

$$\frac{\partial \omega_2^*}{\partial \underline{\theta}} = \begin{cases} 1 & \text{if } \underline{\theta} \leq \mu_1, \\ \frac{\partial \omega_1^*}{\partial \underline{\theta}} (1 - F_2^P(1 + \omega_1^* - \underline{\theta})) + F_2^P(1 + \omega_1^* - \underline{\theta}) & \text{if } \underline{\theta} > \mu_1. \end{cases}$$

Thus, $\frac{\partial r^*}{\partial \underline{\theta}} = \frac{1}{2}(\frac{\partial \omega_1^*}{\partial \underline{\theta}} - \frac{\partial \omega_2^*}{\partial \underline{\theta}}) < 0$.

We now establish that r^* is the unique interior local maximum of Π_d . We can write Π_d as

$$\begin{aligned} \Pi_d(r) &= \mathbb{E}_{\boldsymbol{\theta}} \left[\theta_1 Q_1^e(\boldsymbol{\theta}) - \frac{Q_1^e(\boldsymbol{\theta})^2}{2} - \int_{\omega_1}^{\theta_1} Q_1^e(y, \theta_2) dy \right] - \left(\omega_1 r - \frac{r^2}{2} \right) \\ &\quad + \mathbb{E}_{\boldsymbol{\theta}} \left[\theta_2 Q_2^e(\boldsymbol{\theta}) - \frac{Q_2^e(\boldsymbol{\theta})^2}{2} - \int_{\omega_2}^{\theta_2} Q_2^e(\theta_1, y) dy \right] - \left(\omega_2(1-r) - \frac{(1-r)^2}{2} \right), \end{aligned} \quad (\text{B.21})$$

where the worst-off types are treated parametrically. Differentiating with respect to r , we get

$$\frac{\partial \Pi_d(r)}{\partial r} = (q_1^e(\omega_1) - r) \frac{d\omega_1}{dr} - \omega_1 + r + (q_2^e(\omega_2) - 1 + r) \frac{d\omega_2}{dr} + \omega_2 - (1 - r).$$

With extremal worst-off types, we have $\frac{d\omega_1}{dr} = \frac{d\omega_2}{dr} = 0$, and with interior worst-off types, we have $q_1^e(\omega_1) = r$ and $q_2^e(\omega_2) = 1 - r$, which gives us, regardless of whether worst-off types are interior or not, $\frac{\partial \Pi_d(r)}{\partial r} = \omega_2 - \omega_1 + 2r - 1$. Thus, at any interior local optimum, we have

$$r = \frac{1 + \omega_1 - \omega_2}{2}. \quad (\text{B.22})$$

With interior worst-off types, $\frac{\partial^2 \Pi_d(r)}{\partial r^2} = -\frac{\partial \omega_1}{\partial r} + \frac{\partial \omega_2}{\partial r} + 2$, where $\frac{\partial \omega_1}{\partial r}$ is positive and $\frac{\partial \omega_2}{\partial r}$ is negative and either $\frac{\partial \omega_1}{\partial r} = 2$ or $\frac{\partial \omega_2}{\partial r} = -2$, ensuring that Π_d is concave. Thus, there is a local maximum at $r^* = \frac{1 + \omega_1^* - \omega_2^*}{2}$, where the “star” notation now makes explicit that ω_1^* , ω_2^* , and r^* are determined simultaneously.

With extremal worst-off types, $\frac{\partial^2 \Pi_d(r)}{\partial r^2} = 2$, implying that Π_d is convex, and so there can be no local maximum in this case, but there can be a local minimum. We consider the four possibilities in turn: 1. A local minimum when $(\omega_1, \omega_2) = (1, \underline{\theta})$ requires $r = \frac{2-\theta}{2}$ and $1 - r \leq q_2^e(\underline{\theta}) = \frac{1+\theta-\mu_1}{2}$, which requires $\mu_1 \leq 1$, a contradiction; 2. For $(\omega_1, \omega_2) = (1, 1 + \underline{\theta})$, there is a local minimum at $r = 0$ if $\underline{\theta} = 1$, and otherwise no local minimum. 3. A local minimum when $(\omega_1, \omega_2) = (2, 1 + \underline{\theta})$ requires $r = \frac{2-\theta}{2} \geq q_1^e(2) = \frac{3-\mu_2^P-\theta}{2}$, which requires $\mu_2^P \geq 1$, a contradiction. 4. There is a local minimum with $(\omega_1, \omega_2) = (2, \underline{\theta})$ at $r = \frac{3-\theta}{2}$ (because $\mu_1 \leq 2$ and $\mu_2^P \geq 0$ imply that $\frac{3-\theta}{2} \geq q_1^e(2)$ and $\frac{\theta-1}{2} \leq q_2^e(\underline{\theta})$).

This leaves us to consider cases in which one worst-off type is extremal and one is interior. Working through these cases, one can show that we can never have $\frac{\partial \Pi_d(r)}{\partial r} = 0$. Specifically, one can confirm that when one worst-off type is interior and the other is extremal, one cannot have (B.22).

Case 1: $r < q_1^e(1)$ and $1 - r \in (q_2^e(\underline{\theta}), q_2^e(1 + \underline{\theta}))$. Then $\omega_1 = 1$ and ω_2 is defined by $1 - r = q_2^e(\omega_2)$. If $\omega_2 < 2$, then $\omega_2 = 1 - 2r + \mu_1$ and so (B.22) implies $\mu_1 = 1$, which is a contradiction. If $\omega_2 > 2$, then (B.22) implies $r < 0$, which is a contradiction.

Case 2: $r \in (q_1^e(1), q_1^e(2))$ and $1 - r > q_2^e(1 + \underline{\theta})$. Then $\omega_2 = 1 + \underline{\theta}$ and ω_1 is defined by $r = q_1^e(\omega_1)$. If $\omega_1 > \underline{\theta}$, then $\omega_1 = 2r - 1 + \mu_2^P + \underline{\theta}$ and (B.22) implies $\mu_2^P = 1$, which is a contradiction. If $\omega_1 < \underline{\theta}$, then (B.22) implies $r < 0$, which is a contradiction.

Case 3: $r > q_1^e(2)$ and $1 - r \in (q_2^e(\underline{\theta}), q_2^e(1 + \underline{\theta}))$. Then $\omega_1 = 2$ and ω_2 is defined by $1 - r = q_2^e(\omega_2)$. If $\omega_2 < 2$, then $\omega_2 = 1 - 2r + \mu_1$ and (B.22) implies $\mu_1 = 2$. If $\omega_2 > 2$, then $1 - r = \frac{\omega_2 - 1}{2} + \int_{\omega_2 - 1}^2 \frac{1}{2} F_1(x) dx$, which implies that $\omega_2 - 3 + 2r = - \int_{\omega_2 - 1}^2 F_1(x) dx < 0$, while (B.22) implies $\omega_2 - 3 + 2r = 0$, a contradiction.

Case 4: $r \in (q_1^e(1), q_1^e(2))$ and $1 - r < q_2^e(\underline{\theta})$. Then $\omega_2 = \underline{\theta}$ and ω_1 is defined by $r = q_1^e(\omega_1)$. If $\omega_1 > \underline{\theta}$, then $\omega_1 = 2r - 1 + \mu_2^P + \underline{\theta}$ and (B.22) implies $\mu_2^P = 0$, a contradiction. If $\omega_1^e < \underline{\theta}$, then $r = \frac{1}{2} \int_1^{2 + \omega_1 - \underline{\theta}} F_2^P(x) dx$, while (B.22) implies that $2r = 1 + \omega_1 - \underline{\theta}$, so we have $1 + \omega_1 - \underline{\theta} = \int_1^{2 + \omega_1 - \underline{\theta}} F_2^P(x) dx < (1 + \omega_1 - \underline{\theta}) F_2^P(2 + \omega_1 - \underline{\theta}) < 1 + \omega_1 - \underline{\theta}$, a contradiction.

It remains to consider the possibility of a local maximum at the boundary. First consider $r = 0$. In this case, $\omega_1 = 1$ and $\omega_2 = 1 + \underline{\theta}$ and so we have $\left. \frac{\partial \Pi_d(r)}{\partial r} \right|_{r=0} = \underline{\theta} - 1 \geq 0$, which contradicts a local maximum at $r = 0$ when $\underline{\theta} > 1$. Turning to $r = 1$, in this case, $\omega_1 = 2$ and $\omega_2 = \underline{\theta}$ and so we have $\left. \frac{\partial \Pi_d(r)}{\partial r} \right|_{r=1} = \underline{\theta} - 1 \geq 0$, implying that there is a local maximum at $r = 1$ for $\underline{\theta} > 1$.

It follows that when $\underline{\theta} = 1$, Π_d is quasiconcave with a unique maximum at $r^* = \frac{1 + \omega_1^* - \omega_2^*}{2} = \frac{1 + \mu_1 - \mu_2^P}{2}$, and that when $\underline{\theta} \in (1, 3)$, Π_d is increasing at $r = 0$, has a unique interior local maximum at $r^* = \frac{1 + \omega_1^* - \omega_2^*}{2}$, has a unique interior local minimum at $r = \frac{3 - \underline{\theta}}{2}$, and is increasing at $r = 1$. $\square$

Lemma B.14 establishes that for $\underline{\theta} \in (1, 2)$, $\mathcal{E}_d(\underline{\theta})$ has the form described in the statement of the theorem, but it remains to establish parts (i) and (ii) of the theorem.

Together with Lemma B.14, the next lemma implies that for $w = 1/2$ and $\underline{\theta} = 1$, ex post efficiency is possible if and only if $r = \frac{1 + \mu_1^P - \mu_2^P}{2} = \frac{1 + \mu_1 - \mu_2}{2}$.

Lemma B.15. *If $\underline{\theta} = 1$, then $\Pi_d(\frac{1 + \mu_1^P - \mu_2^P}{2}) = 0$.*

Proof. For $\underline{\theta} = 1$, we have $Q_i^e(\underline{\theta}) = \frac{1 + \theta_i - \theta_j}{2}$, and so it is straightforward to show using (B.21)

that

$$\begin{aligned}\Pi_d(r) &= \frac{1}{2}\mu_1\mu_2 + \frac{1}{2}(\omega_1(1-\mu_2) + \omega_2(1-\mu_1)) + \frac{1}{4}(\omega_1^2 + \omega_2^2) \\ &\quad - (\omega_1 r + \omega_2(1-r) - r^2/2 - (1-r)^2/2) - \frac{1}{4}.\end{aligned}$$

It then follows that $\Pi_d(r)$ is strictly concave at its maximizer $r = \frac{1+\mu_1-\mu_2}{2}$ and that $\Pi_d(\frac{1+\mu_1-\mu_2}{2}) = 0$. The result follows noting that for $\underline{\theta} = 1$, we have $\frac{1+\mu_1-\mu_2}{2} = \frac{1+\mu_1^P-\mu_2^P}{2}$. $\square$

Together with Lemma B.14, the next lemma implies that for $w = 1/2$ and $\underline{\theta} \in (1, \min\{\mu_1, 2 - \mu_2^P\}]$, ex post efficiency is possible for all $r \in [\underline{r}, \bar{r}]$ with $0 \leq \underline{r} < r^* < \bar{r} \leq 1$. The lemma shows that $\Pi_d(r^*) > 0$, and then by the continuity of $\Pi_d(r)$, it follows that $\Pi_d(r) > 0$ for r in an open interval around r^* .

Lemma B.16. *For all $\underline{\theta} \in (1, \min\{\mu_1, 2 - \mu_2^P\}]$, $\Pi_d(r^*) > 0$.*

Proof of Lemma B.16. From Lemma B.15, if $\underline{\theta} = 1$, then $\Pi_d(r^*) = 0$. We show that $\Pi(r^*)$ is increasing in $\underline{\theta}$ for $\underline{\theta} \in (1, \min\{\mu_1, 3 - \mu_2^P\}]$. For $\underline{\theta} \in [1, 2]$, we have

$$\begin{aligned}\frac{\partial \Pi_d(r^*)}{\partial \underline{\theta}} &= \frac{1}{2}(F_2(\min\{1 + \omega_1^*, 1 + \underline{\theta}\})(2 - \omega_1^*) - (2 - \mu_1)) \\ &\quad + \frac{1}{2}\left(\int_1^{\underline{\theta}} F_1(\theta_1)(1 - F_2(1 + \theta_1))d\theta_1 + \int_{\min\{1+\omega_1^*, 1+\underline{\theta}\}}^{1+\underline{\theta}} (3 - \theta_2)dF_2(\theta_2)\right).\end{aligned}$$

As shown in Lemma B.14, if $\underline{\theta} \in [1, \min\{\mu_1, 2 - \mu_2^P\}]$, then $\omega_1^* = \mu_1$, and so in this case,

$$\frac{\partial \Pi_d(r^*)}{\partial \underline{\theta}} = \frac{1}{2}\int_1^{\underline{\theta}} F_1(\theta_1)(1 - F_2(1 + \theta_1))d\theta_1 \geq 0,$$

with a strict inequality if $\underline{\theta} > 1$, establishing that $\Pi_d(r^*)$ is increasing in $\underline{\theta}$ for $\underline{\theta} \in (1, \min\{\mu_1, 2 - \mu_2^P\}]$, and so $\Pi_d(r^*) > 0$ for $\underline{\theta} \in (1, \min\{\mu_1, 2 - \mu_2^P\}]$. $\square$

The next lemma shows that for any $\underline{\theta} \in [1, 3)$, $0 \notin \mathcal{R}_d(\underline{\theta})$, with the implication that for $w = 1/2$ and $\underline{\theta} \in (1, \min\{\mu_1, 2 - \mu_2^P\}]$, ex post efficiency is possible for all $r \in [\underline{r}, \bar{r}]$ with $0 < \underline{r} < r^* < \bar{r} \leq 1$.

Lemma B.17. *For all $\underline{\theta} \in [1, 3)$, $\Pi_d(0) < 0$.*

Proof. Define welfare $W(\boldsymbol{\theta}) \equiv \sum_{i=1}^2 \left(\theta_i Q_i^e(\boldsymbol{\theta}) - \frac{Q_i^e(\boldsymbol{\theta})^2}{2} \right)$ and let W_i denote the partial derivative of W with respect to its i -th argument. By the envelope theorem, $W_i(\boldsymbol{\theta}) = Q_i^e(\boldsymbol{\theta})$. In the VCG mechanism, expected surplus under binding IR for agents' worst-off types is $\mathbb{E}_{\boldsymbol{\theta}} [W(\omega_1, \theta_2) + W(\theta_1, \omega_2) - W(\theta_1, \theta_2)] - O(r)$, where $O(r) \equiv \omega_1 r + \omega_2(1-r) - r^2/2 - (1-r)^2/2$.

$r)^2/2$. Thus, we have

$$\begin{aligned}
\Pi_d(r) &= \mathbb{E}_{\theta} [W(\omega_1, \theta_2) + W(\theta_1, \omega_2) - W(\theta_1, \theta_2)] - O(r) \\
&= \int_{\underline{\theta}}^{1+\underline{\theta}} \int_1^2 (W(\omega_1, \theta_2) + W(\theta_1, \omega_2) - W(\theta_1, \theta_2)) dF_1(\theta_1) dF_2(\theta_2) - O(r) \\
&= W(\omega_1, 1 + \underline{\theta}) + W(2, \omega_2) - W(2, 1 + \underline{\theta}) - \int_1^2 (Q_1^e(\theta_1, \omega_2) - Q_1^e(\theta_1, 1 + \underline{\theta})) F_1(\theta_1) d\theta_1 \\
&\quad - \int_{\underline{\theta}}^{1+\underline{\theta}} (Q_2(\omega_1, \theta_2) - q_2^e(\theta_2)) F_2(\theta_2) d\theta_2 - O(r),
\end{aligned}$$

where the final equality follows by repeated integration by parts. For $r = 0$, we have $(\omega_1, \omega_2) = (1, 1 + \underline{\theta})$ and so $O(0) = \underline{\theta} + 1/2 = W(1, 1 + \underline{\theta})$, implying that for $\underline{\theta} \in [1, 3)$, $\Pi_d(0) = - \int_{\underline{\theta}}^{1+\underline{\theta}} [Q_2(1, \theta_2) - q_2^e(\theta_2)] F_2(\theta_2) d\theta_2 < 0$, where the inequality uses that $Q_2(1, \theta_2) \geq q_2(\theta_2)$, with a strict inequality for $\underline{\theta}$ in an open set of $\underline{\theta} \in [1, 3)$. $\square$

Now we turn to $r^{top}(\underline{\theta})$. Recall from the above lemmas that for $\underline{\theta} = 1$, $\Pi_d(r)$ is quasi-concave with unique interior maximizer r^* with $\Pi_d(r^*) = 0$. Thus, for $\underline{\theta} = 1$, $\Pi_d(1) < 0$. In the lemma below, we show that there exists $\bar{\tau} \in (1, 2)$ such that for $\underline{\theta} = \bar{\tau}$, $\Pi_d(1) = 0$. Further, because for $\underline{\theta} \in (1, 3)$, $\Pi_d(1)$ is increasing in $\underline{\theta}$, we have $\Pi_d(1) > 0$ for all $\underline{\theta} \in (\bar{\tau}, 3)$. From this, it follows that for $\underline{\theta} \in [\bar{\tau}, 3)$, there exists $r^{top}(\underline{\theta})$ such that $\Pi_d(r^{top}(\underline{\theta})) = 0$, with $r^{top}(\bar{\tau}) = 1$ and for $\underline{\theta} \in (\bar{\tau}, 3)$, $r^{top}(\underline{\theta}) \in (0, 1)$, where the result that $r^{top}(\underline{\theta}) > 0$ follows from Lemma B.17. Indeed, because $\Pi_d(r)$ has a local minimum at $r = \frac{3-\underline{\theta}}{2}$, if $\Pi_d(\frac{3-\underline{\theta}}{2}) \leq 0$, then $r^{top}(\underline{\theta}) \geq \frac{3-\underline{\theta}}{2}$, and if $\Pi_d(\frac{3-\underline{\theta}}{2}) > 0$, then we can let $r^{top}(\underline{\theta}) = r^*$.

Lemma B.18. *There exists $\bar{\tau} \in (1, 2)$ such that for all $\underline{\theta} \in [\bar{\tau}, 3)$, $\Pi_d(1) \geq 0$.*

Proof. Because $Q_1^e(\theta_1, \theta_2) \geq 0$, we have $\int_{\theta_1}^2 Q_1^e(y, \theta_2) dy \geq 0$, and because $Q_2^e(\theta_1, \theta_2) \in [0, 1]$ and $\omega_2 \in [\underline{\theta}, 1 + \underline{\theta}]$, we have $\int_{\omega_2}^{\theta_2} Q_2^e(\theta_1, y) dy \leq \int_{\underline{\theta}}^{\theta_2} Q_2^e(\theta_1, y) dy \leq \theta_2 - \underline{\theta}$. Also, these inequalities hold strictly for all types in an open subset of $[1, 2] \times [\underline{\theta}, 1 + \underline{\theta}]$ as long as $\underline{\theta} < 3$. Consequently,

$$\mathbb{E}_{\theta} \left[\sum_{i \in \mathcal{N}} \left(\theta_i Q_i^e(\theta) - \frac{Q_i^e(\theta)^2}{2} \right) + \int_{\theta_1}^2 Q_1^e(y, \theta_2) dy - \int_{\omega_2}^{\theta_2} Q_2^e(\theta_1, y) dy \right] \geq \mathbb{E}_{\theta} \left[\sum_{i \in \mathcal{N}} \left(\theta_i Q_i^e(\theta) - \frac{Q_i^e(\theta)^2}{2} \right) - (\theta_2 - \underline{\theta}) \right],$$

with a strict inequality for $\underline{\theta} < 3$. Moreover, $\mathbb{E}_{\theta} \left[\sum_{i \in \mathcal{N}} \left(\theta_i Q_i^e(\theta) - \frac{Q_i^e(\theta)^2}{2} \right) \right] \geq \mathbb{E}_{\theta} [\theta_2 - \frac{1}{2}]$, which is simply saying that expected social surplus must be at least as large as the social surplus obtained when always allocating everything to agent 2. Using the result that for $r \geq \frac{3-\underline{\theta}}{2}$, we have $r \geq q_1^e(2)$, and so $\omega_1 = 2$, it follows that for $\underline{\theta} \in (1, 3)$:

$$\begin{aligned}
\Pi_d(1) &= \mathbb{E}_{\theta} \left[\sum_{i \in \mathcal{N}} \left(\theta_i Q_i^e(\theta) - \frac{Q_i^e(\theta)^2}{2} \right) + \int_{\theta_1}^2 Q_1^e(y, \theta_2) dy - \int_{\omega_2}^{\theta_2} Q_2^e(\theta_1, y) dy \right] - \frac{3}{2} \\
&> \mathbb{E}_{\theta} \left[\theta_2 - \frac{1}{2} - (\theta_2 - \underline{\theta}) \right] - \frac{3}{2} = \underline{\theta} - 2 \geq 0.
\end{aligned}$$

Further, differentiating $\Pi_d(1)$ with respect to $\underline{\theta}$, we obtain $\frac{\partial \Pi_d(1)}{\partial \underline{\theta}} = \mathbb{E}_{\theta_1}[Q_2^e(\theta_1, \underline{\theta})] > 0$. Because $\Pi_d(1) = 0$ when $\underline{\theta} = 1$ and $\Pi_d(1) > 0$ when $\underline{\theta} = 2$, and because $\Pi_d(1)$ is increasing in $\underline{\theta}$, it follows by continuity that there exists a unique $\bar{\tau} < 2$ such that for all $\underline{\theta} \in [\bar{\tau}, 3)$, $\Pi_d(1) \geq 0$ holds, completing the proof. $\square$

The result that for $\underline{\theta} \in [1, 3)$ ex post efficiency requires $w = 1/2$ follows because, as long as $\underline{\theta} < 3$, there exist θ_1 and θ_2 such that $\frac{1+\theta_1-\theta_2}{2} \in (0, 1)$, and ex post efficiency requires that for all such θ_1 and θ_2 , $1 + \theta_1 - \theta_2 = 1 + \Psi_{1, \frac{w}{\rho}}(\theta_1, \hat{\theta}_1) - \Psi_{2, \frac{1-w}{\rho}}(\theta_2, \hat{\theta}_2)$. In particular, ex post efficiency and the bargaining mechanism require that the ranking of the actual types and the weighted virtual types be the same. In an adaptation of Lemma B.10 for the decreasing marginal value case, given an open set of types for agent 1 and agent 2 such that the ranking of $1 + \theta_1$ and θ_2 varies on that set, the ranking of $1 + \Psi_{1, \frac{w}{\rho}}(\theta_1, \omega_1)$ and $\Psi_{2, \frac{1-w}{\rho}}(\theta_2, \omega_2)$ cannot always match that of $1 + \theta_1$ and θ_2 if the agents' bargaining weights differ, establishing the result. The proof of the remaining results on bargaining weights are analogous to those for constant marginal values and so are omitted. This completes the proof of Theorem 3. $\blacksquare$

C Intuition based on implementation

In Appendix C.1, we provide intuition for Proposition 1 based on optimal take-it-or-leave-it offers. As discussed there, for the case of two agents, where one has all the bargaining weight, the bargaining mechanism corresponds to a Texas shootout in which the agent with all the bargaining weight makes a take-it-or-leave-it offer p , with the interpretation that the other agent can choose whether to buy agent 1's r units, paying rp , or sell its $1 - r$ units, receiving $(1 - r)p$.

In Appendix C.2, we provide intuition for the bargaining mechanism with decreasing marginal values based on optimal nonlinear tariffs.

C.1 Optimal take-it-or-leave-it offers with constant marginal values

To provide intuition for Proposition 1, consider the case with $n = 2$, $w = 1$, and $r \in (0, 1)$ and assume, first, that player 1 simply makes a take-it-or-leave-it offer p , with the interpretation that 2 can choose whether to buy 1's share r , paying rp , or sell its share $1 - r$, receiving $(1 - r)p$. For this illustration, assume that Ψ_i^S and Ψ_i^B are increasing. The optimization problem for 1 is

$$\max_{p \in [\underline{\theta}, \theta + 1]} (1 - r)(\theta_1 - p)F_2(p) + r(p - \theta_1)(1 - F_2(p)),$$

whose maximizer $p^*(\theta_1, r)$ satisfies

$$\theta_1 = (1 - r)\Psi_2^S(p^*(\theta_1, r)) + r\Psi_2^B(p^*(\theta_1, r)). \quad (\text{C.23})$$

This is intuitive because it corresponds to the convex combination of the optimal take-it-or-leave-it offers if agent 1 is a buyer ($r = 0$) and if agent 1 is a seller ($r = 1$). Because we assume for this illustration that Ψ_2^S and Ψ_2^B are increasing, the right side of (C.23) is increasing in p^* and $p^*(\theta_1, r)$ is well defined. If F_2 is the uniform distribution on $[0, 1]$, then $p^*(\theta_1, r) = \frac{\theta_1 + r}{2}$.³ Even though take-it-or-leave-it offers are optimal for $r \in \{0, 1\}$, they are, by Proposition 1, not optimal for $r \in (0, 1)$. The issue is that these leave too much money on the table.

To see this, assume that F_1 and F_2 are both uniform on $[0, 1]$. The expected utility of agent 2's worst-off type, ω_2 , given the pricing rule $p^*(\theta_1, r)$, is $(1 - r)\omega_2 - \int_0^{1-r} p^*(\theta_1, r)d\theta_1 + (1 - r) \int_{1-r}^1 p^*(\theta_1, r)d\theta_1 = (1 - r)\omega_2 + \frac{1}{4}r(1 - r)$ because with probability $1 - r$, agent 2 gets to consume the good when its type is ω_2 . Because $(1 - r)\omega_2$ is the value of agent 2's outside option when its type is ω_2 , it follows that, given the pricing rule $p^*(\theta_1, r)$, agent 2 can be charged the tax $\frac{1}{4}r(1 - r)$ without violating its IC and IR constraints. Naturally, if the price is lower than $p^*(\theta, r)$ when $\theta_1 < 1 - r$ and larger than that for $\theta_1 > 1 - r$, then agent 2's utility when of type ω_2 increases, which means that the tax can be increased. Supposing that agent 1 sets the price $p^*(\theta_1, r - \Delta_S)$ for $\theta_1 < 1 - r$ and $p^*(\theta_1, r + \Delta_B)$ for $\theta_1 > 1 - r$, where $\Delta_S, \Delta_B \geq 0$, the tax becomes $T(\Delta_S, \Delta_B) = \frac{1}{4}r(1 - r)(1 + 2(\Delta_S + \Delta_B))$, which is increasing in Δ_S and Δ_B . So the question is how large Δ_S and Δ_B should be. When agent 1's type is $\theta_1 < 1 - r$, maximizing agent 1's payoff

$$\begin{aligned} & (1 - r)(\theta_1 - p^*(\theta_1, r - \Delta_S))F_2(p^*(\theta_1, r - \Delta_S)) \\ & + r(p^*(\theta_1, r - \Delta_S) - \theta_1)(1 - F_2(p^*(\theta_1, r - \Delta_S))) + T(\Delta_S, \Delta_B) \end{aligned} \quad (\text{C.24})$$

over Δ_S yields the maximizer $\Delta_S = r(1 - r)$, while maximizing

$$\begin{aligned} & (1 - r)(\theta_1 - p^*(\theta_1, r + \Delta_B))F_2(p^*(\theta_1, r + \Delta_B)) \\ & + r(p^*(\theta_1, r + \Delta_B) - \theta_1)(1 - F_2(p^*(\theta_1, r + \Delta_B))) + T(\Delta_S, \Delta_B) \end{aligned} \quad (\text{C.25})$$

over Δ_B yields the maximizer $\Delta_B = r(1 - r)$. The resulting pricing rule does not correspond to the optimal allocation rule.⁴ However, the maximizers of the ex ante expected utility of

³It is interesting to note that $p^*(\theta_1, r)$ would be the optimal price set by agent 1 in a *Texas shootout*, which is a mechanism for dissolving a partnership, whereby "one owner names a price and the other owner is compelled to either purchase the first owner's shares or sell his own shares at the named price" (Brooks et al., 2010, p. 649).

⁴The issue is that maximizing the expected payoff for a given type θ_1 below (above) $1 - r$ neglects the externality that further price reductions (increases) have on all other types. Maximizing ex ante expected utility accounts for these externalities and results in the optimal allocation rule. Put differently, if every type

1 over Δ_S and Δ_B are indeed $\Delta_S^* = r$ and $\Delta_B^* = 1 - r$, resulting in the optimal allocation rule. Observe also that $T(\Delta_S^*, \Delta_B^*) = \frac{3}{4}r(1 - r)$,

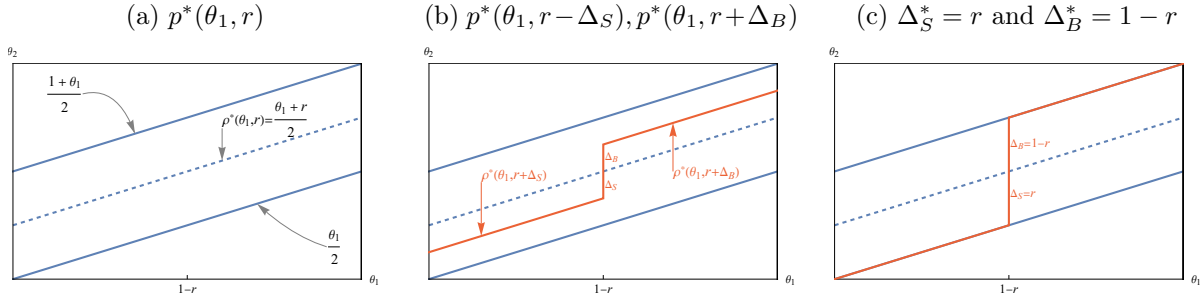


Figure C.2: Take-it-or-leave-it offers. By choosing $\Delta_S > 0$ and $\Delta_B > 0$, the utility of the worst-off type and therefore the tax that can be extracted increases. The optimal choices are $\Delta_S^* = r$ and $\Delta_B^* = 1 - r$.

C.2 Optimal nonlinear tariffs with decreasing marginal values

As noted, with constant marginal values and extremal ownership, the seller-optimal and buyer-optimal bilateral bargaining mechanisms can be implemented with a take-it-or-leave-it price offer from the agent with all the bargaining weight. We now show that this property extends to the model with decreasing marginal values, where the agent with all the bargaining weight can offer any nonlinear tariff it wishes. For the purposes of this subsection, we assume that Ψ_i^S and Ψ_i^B are increasing, and so invertible.

We begin with the case with $r = 1$ and $w = 1$, in which case the seller (agent 1) of type θ_1 offers the seller-optimal tariff $t^S(q; \theta_1)$ to agent 2. Agent 2 then chooses the quantity $q \in [0, 1]$ that it will consume and pays $t^S(q; \theta_1)$ to agent 1, where $t^S(q; \theta_1)$ is given by

$$t^S(q; \theta_1) \equiv \int_0^q (\Psi_2^{B-1}(2y + \theta_1 - 1) - y) dy.$$

In this setting, $Q_1^S(\theta_1, \theta_2) \equiv \min\{1, \max\{0, \frac{1}{2}(1 + \theta_1 - \Psi_2^B(\theta_2))\}\}$ is the seller-optimal quantity for agent 1 when its type is θ_1 and agent 2's type is θ_2 , and, given $t^S(q; \theta_1)$, agent 2 optimally chooses $q = 1 - Q_1^S(\theta_1, \theta_2)$. To see this, note that agent 2 maximizes $q\theta_2 - \frac{1}{2}q^2 - t(q; \theta_1)$ over $q \in [0, 1]$. The first-order condition for an interior solution is $q = \frac{1}{2}(1 - \theta_1 + \Psi_2^B(\theta_2))$, which is equal to $1 - Q_1^S(\theta_1, \theta_2)$ when it is interior. Otherwise, agent 2 chooses the corresponding boundary quantity. Thus, the buyer's optimal choice of q given t^S implements the seller-optimal allocation.

$\theta_1 < 1 - r$ recognizes that all other types less than $1 - r$ decrease their prices by Δ_S , this is as if the impact on the tax increased to $\Delta_S/(1 - r)$, and analogously for types larger than $1 - r$, their impact increases to Δ_B/r . Indeed, if in (C.24) and (C.25) the expression $T(\Delta_S, \Delta_B)$ is replaced by $T(\Delta_S/(1 - r), \Delta_B/r)$, then the respective maximizers become $\Delta_S^* = r$ and $\Delta_B^* = 1 - r$.

Figure C.3 illustrates the payment schedule for $F_2(\theta) = (\theta - 1)^2$ on $[1, 2]$. The nonlinear tariff is concave in q , that is, it involves quantity discounts.

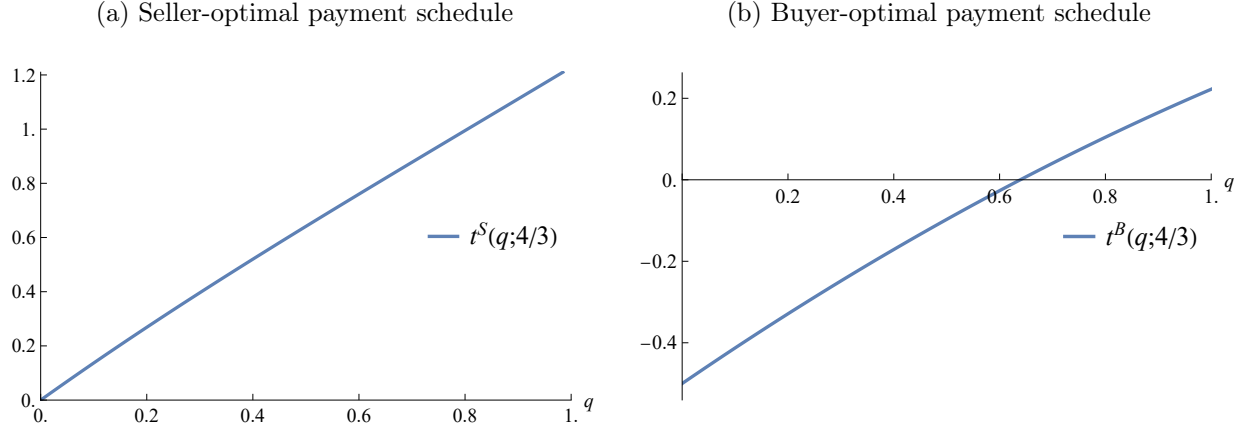


Figure C.3: The seller-optimal and buyer-optimal payment schedules assuming $F_1(\theta) = F_2(\theta) = (\theta - 1)^2$ on $[1, 2]$.

Analogously, when $r = 1$ and $w = 0$, the buyer-optimal mechanism can be implemented through a nonlinear tariff $t^B(q; \theta_2)$ that agent 2 (the buyer) with type θ_2 offers to agent 1, with the understanding that agent 1 chooses the quantity q to sell to agent 2 in exchange for the payment $t^B(q; \theta_2)$. The buyer-optimal allocation has $Q_2^B(\theta_1, \theta_2) \equiv \min\{1, \max\{0, \frac{1}{2}(1 - \Psi_1^S(\theta_1) + \theta_2)\}\}$, which is implemented with the following payment schedule, which is the amount paid by agent 2 to agent 1 as a function of the quantity q that agent 1 chooses to sell to agent 2:

$$t^B(q; \theta_2) \equiv K(\theta_2) - \int_0^{1-q} (\Psi_1^{S-1}(2x - 1 + \theta_2) - x) dx,$$

where $K(\theta_2)$ is defined below. Given this schedule, agent 1 chooses the quantity q to sell (and the quantity $1 - q$ to retain) to maximize $\theta_1(1 - q) - 1/2(1 - q)^2 + t^B(q; \theta_2)$, whose first-order condition is $q = \frac{1}{2}(1 - \Psi_1^S(\theta_1) + \theta_2)$, ensuring the buyer-optimal allocation for agent 2 of $Q_2^B(\theta_1, \theta_2)$. Defining $K(\theta_2)$ so that agent 1's worst-off type, i.e., $\theta_1 = \bar{\theta}_1$, has payoff $\bar{\theta}_1 - \frac{1}{2}\bar{\theta}_1^2$, which is agent 1's outside option, completes the specification of the payment schedule.⁵ The payment schedule is illustrated in Figure C.3(b) for the case of $F_1(\theta) = (\theta - 1)^2$ on $[1, 2]$. As shown there, agent 1 is paid by agent 2 when it sells a sufficiently large quantity to agent 2, but would have to pay agent 2 if it chose to retain its units for itself.

⁵The payment $K(\theta_2)$ is defined by $Q_1^B(\bar{\theta}_1, \theta_2) - 1/2Q_1^B(\bar{\theta}_1, \theta_2)^2 + K(\theta_2) - \int_0^{Q_1^B(\bar{\theta}_1, \theta_2)} (\Psi_1^{S-1}(2x - 1 + \theta_2) - x) dx = \bar{\theta}_1 - \frac{1}{2}\bar{\theta}_1^2$.

D Fee-setting mechanism details

D.1 Fee-setting for bilateral trade

Consider implementation of the bilateral bargaining mechanism for the case with general $w \in [0, 1]$ and $r \in [0, 1]$. We assume that Ψ_i^S and Ψ_i^B are increasing and so invertible. We consider a fee-setting mechanism in which agent 1 sets the price p and then agent 2 chooses whether to buy agent 1's r units at price p or sell its own $1 - r$ units to agent 1 at price p , as in a Texas shootout. If agent 2 buys r units from agent 1, then agent 2 pays rp to agent 1 and agent 1 pays fee $r\phi^S(p)$ to the intermediary. If agent 2 sells $1 - r$ units to agent 1, then agent 2 receives $(1 - r)p$ from agent 1 and agent 1 pays fee $(1 - r)\phi^B(p)$ to the intermediary. In either case, the intermediary pays fixed amount X_1 to agent 1 and X_2 to agent 2. Thus, agent 1 chooses the price to solve $\max_p \tilde{u}_1(\theta_1, p)$, where $\tilde{u}_1(\theta_1, p)$ is agent 1's interim expected net payoff given price p , ignoring fixed payments:

$$\tilde{u}_1(\theta_1, p) \equiv r(p - \phi^S(p))(1 - F_2(p)) + (\theta_1 - (1 - r)(p + \phi^B(p))) F_2(p) - r\theta_1.$$

Agent 1's maximization problem has first-order condition

$$\begin{aligned} 0 &= r(1 - \phi^{S'}(p))(1 - F_2(p)) - (1 - r)(1 + \phi^{B'}(p))F_2(p) \\ &\quad + [-r(p - \phi^S(p)) + \theta_1 - (1 - r)(p + \phi^B(p))] f_2(p). \end{aligned}$$

To implement the allocation rule of the bargaining mechanism, we need agent 1's selected price, $p(\theta_1)$, to map out the price line for the bargaining problem. To ease notation somewhat, define $\alpha_1 \equiv w/\rho$ and $\alpha_2 \equiv (1 - w)/\rho$, where ρ is the Lagrange multiplier on the no-deficit constraint. Then the price line is defined by $\bar{\Psi}_{1,\alpha_1}(\theta_1, \hat{\theta}_1) = \bar{\Psi}_{2,\alpha_2}(\theta_2, \hat{\theta}_2)$, where $\hat{\theta}_i$ is agent i 's worst-off type. To write out the price line, let z_1 and z_2 be the ironing parameters for agents 1 and 2, respectively, and redefine the inverse weighted virtual cost and inverse weighted virtual value functions to have domain equal to the whole real line. That is, in an abuse of notation, for $J \in \{S, B\}$ and $i \in \{1, 2\}$, define

$$\Psi_{i,\alpha_i}^{J^{-1}}(\theta) \equiv \begin{cases} \underline{\theta}_i & \text{if } \theta < \Psi_{i,\alpha_i}^J(\underline{\theta}_i), \\ \Psi_{i,\alpha_i}^{J^{-1}}(\theta) & \text{if } \theta \in [\Psi_{i,\alpha_i}^J(\underline{\theta}_i), \Psi_{i,\alpha_i}^J(\bar{\theta}_i)], \\ \bar{\theta}_i & \text{if } \theta > \Psi_{i,\alpha_i}^J(\bar{\theta}_i). \end{cases}$$

The price line is illustrated in Figure D.4 and defined by (ignoring ties, which occur with

probability zero by Proposition 1): for $\theta_1 \in [\underline{\theta}_1, \bar{\theta}_1]$ and $z_1 \geq z_2$,

$$\tilde{p}(\theta_1) = \begin{cases} \Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(\theta_1)) & \text{if } \theta_1 < \Psi_{1,\alpha_1}^{S^{-1}}(z_2), \\ \Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^S(\theta_1)) & \text{if } \Psi_{1,\alpha_1}^{S^{-1}}(z_2) < \theta_1 < \Psi_{1,\alpha_1}^{S^{-1}}(z_1), \\ \Psi_{2,\alpha_2}^{S^{-1}}(z_1) & \text{if } \Psi_{1,\alpha_1}^{S^{-1}}(z_1) \leq \theta_1 \leq \Psi_{1,\alpha_1}^{B^{-1}}(z_1), \\ \Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(\theta_1)) & \text{if } \Psi_{1,\alpha_1}^{B^{-1}}(z_1) < \theta_1, \end{cases}$$

and for $z_1 < z_2$,

$$\tilde{p}(\theta_1) = \begin{cases} \Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^S(\theta_1)) & \text{if } \theta_1 < \Psi_{1,\alpha_1}^{S^{-1}}(z_1), \\ \Psi_{2,\alpha_2}^{S^{-1}}(z_1) & \text{if } \Psi_{1,\alpha_1}^{S^{-1}}(z_1) < \theta_1 < \Psi_{1,\alpha_1}^{B^{-1}}(z_1), \\ \Psi_{2,\alpha_2}^{S^{-1}}(\Psi_{1,\alpha_1}^B(\theta_1)) & \text{if } \Psi_{1,\alpha_1}^{B^{-1}}(z_1) \leq \theta_1 \leq \Psi_{1,\alpha_1}^{B^{-1}}(z_2), \\ \Psi_{2,\alpha_2}^{B^{-1}}(\Psi_{1,\alpha_1}^B(\theta_1)) & \text{if } \Psi_{1,\alpha_1}^{B^{-1}}(z_2) < \theta_1. \end{cases}$$

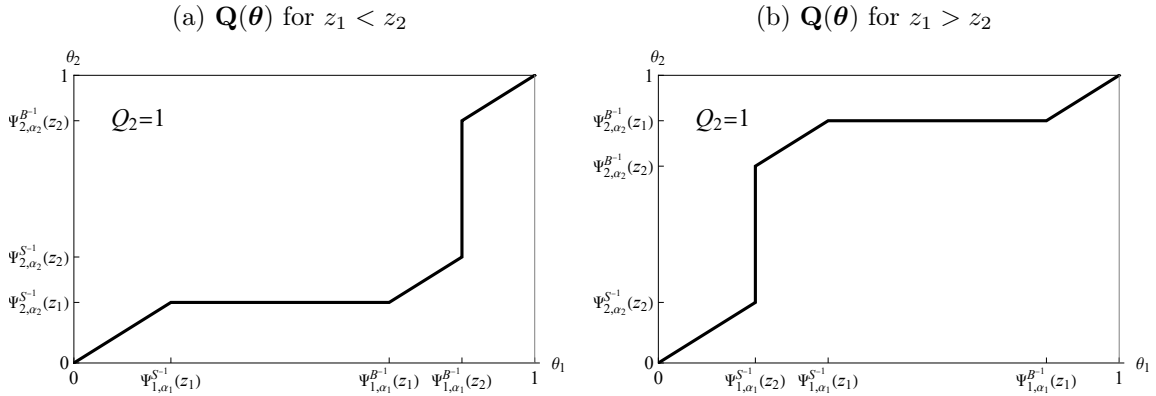


Figure D.4: IIB allocation rule for given z_1 , z_2 , α_1 , and α_2 . Panel (a) assumes $z_1 = 0.4$ and $z_2 = 0.7$, and panel (b) assumes the reverse. Both panels assume that $\alpha_1 = \alpha_2 = 0.1$.

In what follows it will be useful to define the inverse function $\tilde{p}^{-1}(x)$. To ensure that this is defined everywhere, we must address the jump up and flat portion in $\tilde{p}(\theta_1)$. For $z_1 \geq z_2$, the jump occurs at $\theta_1 = \Psi_{1,\alpha_1}^{S^{-1}}(z_2)$ and the flat is at $\Psi_{2,\alpha_2}^{S^{-1}}(z_1)$. For $z_1 < z_2$, the jump occurs at $\theta_1 = \Psi_{1,\alpha_1}^{B^{-1}}(z_2)$ and the flat is at $\Psi_{2,\alpha_2}^{B^{-1}}(z_1)$. So, for $z_1 \geq z_2$ and $x \in (\Psi_{2,\alpha_2}^{S^{-1}}(z_2), \Psi_{2,\alpha_2}^{B^{-1}}(z_2))$, we define $\tilde{p}^{-1}(x) = \Psi_{1,\alpha_1}^{S^{-1}}(z_2)$, and we define $\tilde{p}^{-1}(\Psi_{2,\alpha_2}^{B^{-1}}(z_1)) = \Psi_{1,\alpha_1}^{S^{-1}}(z_1)$; and for $z_1 < z_2$ and $x \in (\Psi_{2,\alpha_2}^{S^{-1}}(z_2), \Psi_{2,\alpha_2}^{B^{-1}}(z_2))$, we define $\tilde{p}^{-1}(x) = \Psi_{1,\alpha_1}^{B^{-1}}(z_2)$, and we define $\tilde{p}^{-1}(\Psi_{2,\alpha_2}^{S^{-1}}(z_1)) = \Psi_{1,\alpha_1}^{S^{-1}}(z_1)$. Note that if $z_1 \geq z_2$, then $\tilde{p}^{-1}(x)$ is not differentiable at $\Psi_{2,\alpha_2}^{B^{-1}}(z_1)$, and for $z_1 < z_2$, it is not differentiable at $\Psi_{2,\alpha_2}^{S^{-1}}(z_1)$. In what follows, we can arbitrarily define the derivatives at these points to be zero.

We now define the fee functions:

$$\phi^S(p) \equiv \frac{1}{r} \int_0^p \tilde{p}^{-1}(x)(1 - F_2(x))dx - \frac{1}{r} \tilde{p}^{-1}(p) + p$$

and

$$\phi^B(p) \equiv \frac{1}{1-r} \int_0^p \tilde{p}^{-1'}(x)(1 - F_2(x))dx - p.$$

Given these fee functions, the first-order condition becomes $(-\tilde{p}^{-1}(p) + \theta_1)f_2(p) = 0$, which is satisfied at $p = \tilde{p}(\theta_1)$, as required, with the second-order condition also satisfied there. Thus, the mechanism induces the allocation of the bargaining mechanism.

Agent 1's interim expected net payoff (ignoring fixed payments) is $u_1(\theta_1) \equiv \tilde{u}_1(\theta_1, \tilde{p}(\theta_1))$. Let $\hat{\theta}_1$ denote the minimizer of this on $[\underline{\theta}_1, \bar{\theta}_1]$. Analogously, agent 2's interim expected net payoff (ignoring fixed payments) is

$$\begin{aligned} u_2(\theta_2) &= \mathbb{E}_{\theta_1}[(1-r)\tilde{p}(\theta_1) \mid \tilde{p}(\theta_1) > \theta_2] \Pr(\tilde{p}(\theta_1) > \theta_2) \\ &\quad + \mathbb{E}_{\theta_1}[\theta_2 - r\tilde{p}(\theta_1) \mid \tilde{p}(\theta_1) < \theta_2] \Pr(\tilde{p}(\theta_1) < \theta_2) - (1-r)\theta_2, \end{aligned}$$

and we denote the minimizer of this on $[\underline{\theta}_2, \bar{\theta}_2]$ by $\hat{\theta}_2$. For example, with uniformly distributed types on $[0, 1]$, we have $u_1(\hat{\theta}_1) < 0$ and $u_2(\hat{\theta}_2) > 0$, which means that for IR to be satisfied, agent 1 must be paid a positive fixed fee, while agent 2 can be taxed with a negative fixed fee.

To pin down the fixed fees, note that the intermediary's expected budget surplus under binding IR is

$$\Pi^I(r) \equiv \mathbb{E}_{\theta_1} [r\phi^S(\tilde{p}(\theta_1))(1 - F_2(\tilde{p}(\theta_1))) + (1-r)\phi^B(\tilde{p}(\theta_1))F_2(\tilde{p}(\theta_1))] + u_1(\hat{\theta}_1) + u_2(\hat{\theta}_2),$$

which, by the definition of ρ , is greater than or equal to zero. Fixed fees are given by $X_1 = -u_1(\hat{\theta}_1) + \mathbf{1}_{w_1 > w_2}S(r) + \mathbf{1}_{w_1 = w_2}S(r)/2$ and $X_2 = -u_2(\hat{\theta}_2) + \mathbf{1}_{w_1 < w_2}S(r) + \mathbf{1}_{w_1 = w_2}S(r)/2$, leaving the intermediary with payoff $\mathbb{E}_{\theta_1} [r\phi^S(\theta_1)(1 - F_2(\theta_1)) + (1-r)\phi^B(\theta_1)F_2(\theta_1)] - X_1 - X_2 = 0$.

This completes the proof of the existence of a fee-setting mechanism that implements the bargaining mechanism.

D.2 Fee-setting for multilateral trade

Going beyond bilateral trade, for the case of one seller and $n - 1$ buyers, an adaptation of the fee-setting mechanism in the online appendix of Loertscher and Marx (2022) applies. We assume that Ψ_i^S and Ψ_i^B are increasing and so invertible. Let agent 1 be the seller, with $r_1 = 1$, and let agents $2, \dots, n$ be buyers with $r_i = 0$ for $i \in \{2, \dots, n\}$. For convenience, assume that all agents have distributions with support $[0, 1]$. The extensive-form game in this case includes an ascending-bid auction, with the following timing: 1. the intermediary announces (and commits to) a *discriminatory clock auction*, which we define below, and fee schedule $\sigma = (\sigma_2, \dots, \sigma_n)$, where σ_i maps the price p paid by the winning buyer i to the seller onto the fee $\sigma_i(p)$ paid by the seller to the intermediary; 2. the seller sets a reserve r

for the auction; 3. the intermediary holds the auction with reserve r , which determines the winning buyer, if any, and the payment to that buyer; 4. given winner i and auction price p , the seller provides the good to buyer i , buyer i pays p to the seller, and the seller pays $\sigma_i(p)$ to the intermediary. If no buyer bids above the reserve, then there is no trade and no payments are made, including no payment to the intermediary.

In the auction portion, we have an ascending clock auction, with the clock price starting at the reserve r and increasing from there. Participants choose when to exit, and when they exit, they become inactive and remain so. The clock stops when only one active bidder remains, with ties broken by randomization. A *discriminatory* clock auction specifies buyer-specific discounts off the final clock price $(\delta_2, \dots, \delta_n)$, where δ_i maps the clock price to buyer i 's discount—activity by buyer i at a clock price of $\hat{p}$ obligates buyer i to pay $\hat{p} - \delta_i(\hat{p})$. By the usual clock auction logic, in the essentially unique equilibrium in non-weakly-dominated strategies, buyer i with type θ_i remains active in the auction until the clock price reaches $\hat{p}$ such that $\hat{p} - \delta_i(\hat{p}) = \theta_i$, and then buyer i exits. We assume that the buyers follow these strategies. The seller chooses the reserve to maximize its expected payoff.

The bargaining mechanism is then implemented by an intermediary that sets auction discounts relative to the clock price $\hat{p}$ of $\delta_i(\hat{p}) \equiv \hat{p} - \Psi_{i, w_i/\rho}^{B^{-1}}(\hat{p})$ and a fee schedule given by, for all $i \in \{2, \dots, n\}$,

$$\sigma_i(p) \equiv p - \Psi_{w_1/\rho}^{S^{-1}}(\Psi_{i, w_i/\rho}^B(\Psi_i^{B^{-1}}(p))).$$

In this case, the seller optimally sets a reserve of $\Psi_{w_1/\rho}^S(\theta_1)$. Then, given our assumption that each buyer i follows its weakly dominant strategy of remaining active until a clock price $\hat{p}$ such that $\hat{p} - \delta_i(\hat{p}) = \theta_i$, buyer i remains active until a price of $\Psi_{i, w_i/\rho}^B(\theta_i)$, and so buyer i wins if and only if

$$\Psi_{i, w_i/\rho}^B(\theta_i) = \max_{j \in \{2, \dots, n\}} \Psi_{j, w_j/\rho}^B(\theta_j) \geq \Psi_{w_1/\rho}^S(\theta_1),$$

which corresponds to the allocation rule of the bargaining mechanism. In equilibrium, if buyer i wins the auction, then the auction ends with a clock price of

$$\hat{p} \equiv \max_{j \in \{2, \dots, n\} \setminus \{i\}} \{\Psi_{w_1/\rho}^S(\theta_1), \Psi_{j, w_j/\rho}^B(\theta_j)\},$$

and the seller is paid $p = \hat{p} - \delta_i(\hat{p})$ by buyer i and the seller pays $\sigma_i(p)$ to the intermediary.

The argument for why a reserve of $\Psi_{w_1/\rho}^S(\theta_1)$ is optimal for the seller is analogous to the argument in Loertscher and Marx (2022). To see this, let $x_i \equiv w_i/\rho$ to reduce notation and define the distribution of buyer i 's weighted virtual type $\Psi_{i, x_i}^B(\theta_i)$ by $\tilde{F}_{i, x_i}(z) = F_i(\Psi_{i, x_i}^{B^{-1}}(z))$, and, letting $\mathbf{x} \equiv (x_1, \dots, x_n)$, define the distribution of the maximum of the weighted virtual

types of buyers other than i by

$$\tilde{F}_{-i,\mathbf{x}}(z) = 1 - \prod_{j \in \mathcal{N} \setminus \{i\}} (1 - \tilde{F}_{j,x_j}(z)).$$

The expected payment to the seller by the buyers given the reserve μ can be written as

$$\begin{aligned} & \sum_{i \in \{2, \dots, n\}} \mathbb{E} \left[\Psi_i^B(\theta_i) \cdot \mathbf{1}_{\Psi_{i,x_i}^B(\theta_i) \geq \max_{j \neq i} \{\mu, \Psi_{j,x_j}^B(\theta_j)\}} \right] \\ &= \sum_{i \in \{2, \dots, n\}} \int_{\min\{1, \Psi_{i,x_i}^{B^{-1}}(\mu)\}}^1 \int_{-\infty}^{\Psi_{i,x_i}^B(\theta_i)} \Psi_i^B(\theta_i) d\tilde{F}_{-i,\mathbf{x}}(z) dF_i(\theta_i) \\ &= \sum_{i \in \{2, \dots, n\}} \int_{\min\{1, \Psi_{i,x_i}^{B^{-1}}(\mu)\}}^1 \Psi_i^B(\theta_i) \tilde{F}_{-i,\mathbf{x}}(\Psi_{i,x_i}^B(\theta_i)) dF_i(\theta_i) \\ &= \sum_{i \in \{2, \dots, n\}} \int_{\min\{1, \Psi_{i,x_i}^B(\Psi_{i,x_i}^{B^{-1}}(\mu))\}}^1 y \tilde{F}_{-i,\mathbf{x}}(\Psi_{i,x_i}^B(\Psi_{i,x_i}^{B^{-1}}(y))) \frac{f_i(\Psi_{i,x_i}^{B^{-1}}(y))}{\Psi_{i,x_i}^{B'}(\Psi_{i,x_i}^{B^{-1}}(y))} dy, \end{aligned}$$

where the final equality uses the change of variables $y = \Psi_i^B(\theta_i)$. Thus, the seller with type θ_1 maximizes its interim expected payoff by choosing μ to solve

$$\max_{\mu} \sum_{i \in \{2, \dots, n\}} \int_{\min\{1, \Psi_{i,x_i}^B(\Psi_{i,x_i}^{B^{-1}}(\mu))\}}^1 (y - \sigma_i(y) - \theta_1) \tilde{F}_{-i,\mathbf{x}}(\Psi_{i,x_i}^B(\Psi_{i,x_i}^{B^{-1}}(y))) \frac{f_i(\Psi_{i,x_i}^{B^{-1}}(y))}{\Psi_{i,x_i}^{B'}(\Psi_{i,x_i}^{B^{-1}}(y))} dy,$$

whose first-order condition when $1 > \Psi_{i,x_i}^B(\Psi_{i,x_i}^{B^{-1}}(\mu))$ is

$$\begin{aligned} 0 &= \sum_{i \in \{2, \dots, n\}} (\Psi_{i,x_i}^B(\Psi_{i,x_i}^{B^{-1}}(\mu)) - \sigma_i(\Psi_{i,x_i}^B(\Psi_{i,x_i}^{B^{-1}}(\mu))) - \theta_1) \tilde{F}_{-i,\mathbf{x}}(\mu) f_i(\Psi_{i,x_i}^{B^{-1}}(\mu)) \Psi_{i,x_i}^{B^{-1}'}(\mu) \\ &= \sum_{i \in \{2, \dots, n\}} (\Psi_{x_1}^{S^{-1}}(\mu) - \theta_1) \tilde{F}_{-i,\mathbf{x}}(\mu) f_i(\Psi_{i,x_i}^{B^{-1}}(\mu)) \Psi_{i,x_i}^{B^{-1}'}(\mu) \end{aligned}$$

where the second equality uses the definition of the fee schedule σ . Given our assumptions, the second-order condition is satisfied when the first-order condition is, and so the seller's problem is solved by $\mu = \Psi_{w_1/\rho}^S(\theta_1)$.

References for the online appendix

- BÖRGERS, T. (2015): *An Introduction to the Theory of Mechanism Design*, Oxford University Press.
- BROOKS, R. R., C. M. LANDEO, AND K. E. SPIER (2010): “Trigger Happy or Gun Shy? Dissolving Common-Value Partnerships with Texas Shootouts,” *RAND Journal of Economics*, 41, 649–673.
- CHEN, Y.-C., W. HE, J. LI, AND Y. SUN (2019): “Equivalence of Stochastic and Deterministic Mechanisms,” *Econometrica*, 87, 1367–1390.
- CLARKE, E. (1971): “Multipart Pricing of Public Goods,” *Public Choice*, 11, 17–33.
- CRAMTON, P., R. GIBBONS, AND P. KLEMPERER (1987): “Dissolving a Partnership Efficiently,” *Econometrica*, 55, 615–632.
- GROVES, T. (1973): “Incentives in Teams,” *Econometrica*, 41, 617–631.
- LOERTSCHER, S. AND L. M. MARX (2022): “Incomplete Information Bargaining with Applications to Mergers, Investment, and Vertical Integration,” *American Economic Review*, 112, 616–649.
- LOERTSCHER, S. AND C. WASSER (2019): “Optimal Structure and Dissolution of Partnerships,” *Theoretical Economics*, 14, 1063–1114.
- MYERSON, R. B. (1981): “Optimal Auction Design,” *Mathematics of Operations Research*, 6, 58–73.
- VICKREY, W. (1961): “Counterspeculation, Auction, and Competitive Sealed Tenders,” *Journal of Finance*, 16, 8–37.