

Bargaining in Vertical Markets

ECON 552/890

Luca Maini*

January 19, 2022

See most recent version [here](#)

CONTENTS

1	Introduction	2
1.1	The original bargaining problem	2
1.2	Introduction to bargaining in vertical markets	4
2	Bargaining in Health Care Markets	6
2.1	A theory of bargaining in vertical markets: Horn and Wolinsky (1988)	6
2.2	An example from health care: Gowrisankaran, Nevo and Town (2015)	9
3	Recent Advances in the Empirical Estimation of Bargaining Models	12
4	Bibliography	15

*These notes are a work in progress. Please check back frequently for an updated version, and email lmaini@email.unc.edu for corrections.

1 INTRODUCTION

1.1 The original bargaining problem

We use the term bilateral bargaining to describe problems where two parties split surplus from trade. Consider two agents, which we will call a seller and a customer. The seller has an object that the customer wants. Each agent has a valuation for the object which we denote as V_s and V_c for the seller and customer respectively. Suppose that $V_s < V_c$, so that there are gains from trade. The question we ask is: at what price will the object be sold?

Before continuing note a couple of things. First, the only economic prediction we can make from this setup is that the transaction should occur. Nothing tells us anything about price, and, in fact, any price $p \in [V_s, V_c]$ is equivalent from the point of view of economic efficiency, since price here is simply a transfer. This means that we have to add more structure in order to obtain a more precise prediction. Second, we have not specified an information structure for this game. Does the seller know the valuation of the customer? What about the other way around? I will largely stay away from these issues and generally assume full information throughout the rest of these notes, but you should know that this is not an inconsequential assumption. It is not uncommon to encounter empirical phenomena that can only be justified by the presence of imperfect or asymmetric information (a very simple example is absence of trade when trade is advantageous). However, allowing for a more flexible information structure generally leads to significant complications in these models, which is why I will generally avoid it.

The **Nash bargaining** solution stipulates that the price at which the product will be traded is the solution to the following maximization problem:

$$\max_p [p - V_s]^{b_s} \times [V_c - p]^{b_c}$$

Here the parameters b_s and b_c indicate the *bargaining power* of the seller and the customer respectively. In most empirical applications b_c and b_s cannot be separately identified so we apply a normalization to the bargaining power parameters and re-write the problem as

$$\max_p [p - V_s]^b \times [V_c - p]^{1-b} \tag{1}$$

This is a monotonic transformation of the problem, so it does not change its predictions. Another commonly applied transformation that is useful for deriving first order conditions is a log transpose, which yields

$$\max_p b \times \log(p - V_s) + (1 - b) \times \log(V_c - p)$$

It is almost always much easier to write down the FOC of the log-transpose of the Nash bargaining problem because we don't have to apply the product rule, and we don't have to deal with the

derivative of power functions.

The solution to this simple problem (which you can check by deriving the FOC yourself) is

$$p = (1 - b) V_s + b V_c$$

In other words, the solution is a linear combination of the valuations of the two parties, with the bargaining weights determining the split.

Some notes

- You can check that the corner cases make sense: if either party has a bargaining power of 1 (i.e. $b = 1$ or $b = 0$), the problem reduces to a simple take-it-or-leave-it offer model where the entire surplus accrues to the party with all the bargaining power.
- The question that this model tries to answer is one of the most basic in economics. It was posed by [Edgeworth \(1881\)](#) as a question of moral philosophy. Edgeworth noted the indeterminacy of the bargaining solution: “It is the interest of both parties that there should be some settlement, [...] but which of these contracts is arbitrary in the absence of arbitration.” Similar notions had actually been around even earlier, in the work of Jevons (*The Theory of Political Economy*, 1871) and Courcelle-Seneuil (*Traité théorique et pratique d’économie politique*, 1858). [Rubinstein \(1982\)](#) provides a good introduction to the problem and derives a microfoundation for the solution as a game of alternating offers where the bargaining parameters are linked to the discount rates of the two parties.
- In empirical applications it is often unclear what “bargaining power” means in practice. In most cases, you should simply treat b as a residual—i.e. an error term that captures variation in the data for which we don’t have a direct explanation. Occasionally however, you may be able to find variation that allows you to identify some properties of b , such as the ability of a given party to consistently obtain “better” deals. We will return to this point later.
- Another important element in empirical applications is the valuation (V_s or V_c). Valuations are often called disagreement or *threat points*, because they represent what either agent can “threaten” to obtain by “disagreeing”, or walking away from the negotiating table. The notion of threat point becomes complicated very quickly in settings with more than two agents. For example, if there are two customers that could purchase the product, the threat point of the seller may not simply be its own valuation, but whatever price the seller can charge to the other customer. That is probably an endogenous value that needs to be pinned down by making some additional assumptions about what would happen if the seller decided to stop negotiating with the first customer. Most empirical research uses very simple assumptions on threat points, and we will do the same here. That being said, this is an area at the frontier of empirical research in the field, and most applications are actually in health care markets. I discuss some recent innovations in the last section of these notes.

- Threat points and bargaining parameters are virtually never directly observed in the data. This is problematic, because while a change in either one will have a similar effect on market outcomes, they have different implications for counterfactuals (namely, bargaining parameters have global effects on all negotiations, whereas threat points only affect a single negotiation). This makes it very important to have a good model of threat points. Unfortunately, this is an area where empirical work is still struggling, despite some recent advances (see, e.g., [Liebman, 2018](#); [Ghili, 2020](#); [Ho and Lee, 2019](#)).

1.2 Introduction to bargaining in vertical markets

Vertical markets are markets where firms contract with one another to supply inputs (e.g. channels and cable providers), establish networks (e.g. consoles and video games), or simply arrange distribution (e.g. wholesalers and retailers). Often, the parties involved in these contracts are large firms that negotiate terms on an individual basis. This introduces two layers of complexity for empirical research. First, there is both great scope for price discrimination and very little actual data on the terms of these contracts, which are usually proprietary. Second, the large number of negotiations that are occurring simultaneously (and the even larger number of negotiations that are *not* occurring) makes it challenging to provide a realistic description of the market that is also amenable to estimation. In the next section, I provide a rigorous introduction to the theory of bargaining in vertical markets. Here, I provide an intuitive presentation based on remarks in [Nevo \(2014\)](#).

Let's start with a simple, concrete example.¹ Consider the problem of Blue Cross Blue Shield of North Carolina (BCBS), which is creating a health plan for patients in the Triangle area, and needs to contract with two large hospital systems: Duke and UNC (assume that all providers in the area are associated with one of those two systems). Also assume that negotiations take the form of Nash Bargaining, and that the outcome of the negotiation is to split the gains equally. Because patients prefer plans that cover more providers, BCBS can earn total profits of \$100 from offering a plan that includes both Duke and UNC. If instead BCBS includes only one hospital system (either Duke or UNC), it earns profits of \$60. How do BCBS, Duke, and UNC split profits in this market?

We can answer this question using the model of bargaining we just discussed in the previous section. Suppose that negotiations occur separately and simultaneously, and that, since all parties stand to gain from putting an agreement to paper, everyone assumes that all negotiations will be successful. Consider the negotiation between BCBS and UNC first. Let p_{UNC} be the transfer from BCBS to UNC. Then, the value of agreeing to a contract for BCBS is $100 - p_{UNC} - p_{DUKE}$ (because we assume that the negotiation with Duke is successful. p_{DUKE} is the transfer to Duke, but it will turn out to not actually matter). The value of disagreeing is $60 - p_{DUKE}$. Conversely, UNC gets p_{UNC} by agreeing, and nothing if negotiations are unsuccessful. The problem we can write down

¹The setting is inspired by an exercise from Allan Collard-Wexler, which you can find [here](#). [Nevo \(2014\)](#) also contains a similar hypothetical problem.

is

$$\begin{aligned} \max_{p_{UNC}} (100 - p_{UNC} - p_{DUKE} - (60 - p_{DUKE}))^{\frac{1}{2}} \times (p_{UNC} - 0)^{\frac{1}{2}} \\ \equiv \max_{p_{UNC}} (40 - p_{UNC})^{\frac{1}{2}} \times (p_{UNC} - 0)^{\frac{1}{2}} \end{aligned}$$

The solution that we derived earlier gives us a simple formula for the equilibrium price:

$$p_{UNC} = \frac{1}{2} \times 40 + \frac{1}{2} \times 0 = 20$$

You can check that the solution to the bargaining problem with Duke is the same.

What would happen if Duke and UNC decided to merge and negotiate as a unit? Instead of two negotiations we have only one. BCBS would make a single payment p_{UD} to both UNC and Duke. The value of agreeing to the contract is $100 - p_{UD}$, and the value of disagreeing is 0. Duke and UNC collectively get p_{UD} by agreeing, and nothing if negotiations are unsuccessful. Following the same script as before we can write the problem as

$$\max_{p_{UD}} (100 - p_{UD})^{\frac{1}{2}} \times (p_{UD} - 0)^{\frac{1}{2}}$$

The solution is

$$p_{UD} = \frac{1}{2} \times 100 + \frac{1}{2} \times 0 = 50$$

The payment received by Duke and UNC when bargaining together is higher than the sum of payments received by the two hospitals when bargaining separately. Why is that? Intuitively this is because the marginal value of an additional hospital decreases with the number of hospitals already available. Another way to say this is that BCBS faces decreasing marginal profits from adding hospitals. Having two hospitals is only slightly better than having two hospitals, but having one hospital is much better than having none. When Duke and UNC negotiate separately, the threat point of BCBS is still pretty good, because they can still agree to terms with the other hospital. When Duke and UNC bargain together the threat point of BCBS becomes much worse and the insurer has to surrender a higher payment.

Some notes

- You could question some of the assumptions in this model (and you'd be at least somewhat right): if BCBS and Duke disagree, wouldn't UNC up its demands? We rule out this possibility because we essentially assume that negotiations are conducted simultaneously and separately from one another. That's an assumption however, which may or may not hold in reality. I discuss these assumptions more in detail in the next section.
- If you change the numbers you can get different results. For example, suppose that the value of having a single hospital in the network is \$50. In other words, the marginal profits

from adding hospitals are flat. In that case, merging does not help UNC and Duke achieve a higher payment. What if having a single hospital is worth only \$40? In that case, merging actually *hurts* Duke and UNC.

- If we had a market with multiple insurers, and we had to analyze the impact of an insurer's merger, the answer would be much more complicated: a single insurer can get a lower price from each hospital because of bargaining gains. This is good, because it means that the insurer costs go down, which lowers premiums and benefits patients buying insurance. However, merging also allows the insurer to charge a higher markup, which increases premiums. The net effect of an insurer merger is therefore ambiguous. We will have more to say about this when we cover mergers later on in the course.

2 BARGAINING IN HEALTH CARE MARKETS

Nash bargaining is the workhorse model to analyze business-to-business interactions, such as supplier-retailer relationships, where each customer is large enough to command a personalized contract—compare this to retail markets, where each customer pays the posted price and there is limited room for individual negotiations.

Most of these business-to-business interactions occur in vertical markets, where an upstream firm supplies a good or service to a downstream firm, which either re-sells it, or uses it as an input in their production functions. Examples from health care markets include insurers negotiating with hospitals to form a network of providers (Ho, 2009; Gowrisankaran et al., 2015; Ho and Lee, 2017), hospitals negotiating with medical device manufacturers (Grennan, 2013; Grennan and Swanson, 2020), and pharmacy benefit managers (or governments) negotiating with pharmaceutical companies (Feng, 2019; Dubois et al., 2019).

2.1 A theory of bargaining in vertical markets: Horn and Wolinsky (1988)

The standard framework for vertical bargaining models comes from Horn and Wolinsky (1988). They consider a model with two symmetric firms (1 and 2) competing in a downstream retail market. Each firm i uses inputs l_i from a single supplier. Input prices are determined through Nash bargaining, and sales in the downstream market occur after negotiation are completed.

We solve using backward induction.

Retail stage

The goal of this stage is to recover $\pi_i(w_i, w_j)$, i.e. the downstream retailer's profits for a given set of input prices w_i and w_j . The two firms play a Cournot-Nash equilibrium where the inverse demand for product i is characterized as

$$p_i(q_i, q_j) = a - cq_j - q_i$$

where c can be either positive (substitute goods) or negative (complementary goods). q_i and q_j are the quantities produced for each of the two goods. We also assume that inputs can be transformed into the final good one-to-one, so that $l_i = q_i$, and that the cost of production is linear. Hence $c(q_i) = w_i q_i$, where w_i is the input price.

The firm's maximization problem can be written as

$$\max_{q_i} p_i(q_i, q_j) \times q_i - w_i q_i$$

and yields equilibrium output

$$q_i^*(w_i, w_j) = l_i(w_i, w_j) = \frac{a(2-c) + cw_j - 2w_i}{4-c^2} \quad (2)$$

and profits

$$\pi_i^*(w_i, w_j) = \left[\frac{a(2-c) + cw_j - 2w_i}{4-c^2} \right]^2 \quad (3)$$

Note that the first complication that appears here is that **the profit function of firm i depends on the bargaining outcome of firm j** . This in turn is going to imply that the structure of bargaining is going to have a big impact on the final equilibrium (e.g. simultaneous vs. sequential bargaining, independent suppliers vs. a monopolistic supplier, vs. something else, etc.). Much of the value of [Horn and Wolinsky \(1988\)](#) lies in carefully unpacking these results.

Bargaining stage

I consider here the situation where there is a single supplier that is bargaining simultaneously with both downstream retailers. This is the most common application that you will find in empirical research, and the most conceptually challenging as well. You can refer to the paper for alternative specifications such as (i) separate, independent suppliers (Section 3), and (ii) a single supplier negotiating sequentially (Section 6).

The supplier's profit function is given by $w_1 l_1 + w_2 l_2$. The definition of a Nash equilibrium in this model is a pair of prices (w_1^S, w_2^S) such that w_i^S is the Nash solution to the bargaining problem between the supplier and firm i and both agents correctly anticipate that w_j^S will be the solution to the other bargaining problem.

The derivation of (w_1^S, w_2^S) is not as straightforward, because it requires specifying the disagreement point. There are many ways to do so, and the right one is not necessarily obvious or simple to implement. Here are a couple of solutions for your consideration.

- The solutions that are most attuned to the concept of a Nash equilibrium involve agents assuming that the outcome of the other negotiation will be unaffected by a disagreement. Hence, the payoff earned by the supplier in case of a disagreement with firm i would simply be $w_j^S l_j$. A natural way to figure out what l_j would be under disagreement is to solve the

downstream competition model for the case of a monopolist, which would essentially yield $q_j^*(w_j^S, \infty)$.

- [Horn and Wolinsky \(1988\)](#) actually do something different, and assume that $l_j = q_i^*(w_i^S, w_j^S)$, i.e. supplier j sticks with the level of output that was optimal under the assumption that both negotiations are completed successfully. They argue that this simplifies exposition, and does not have significantly different implications from the case above. Still, it inserts a degree of internal inconsistency in the model. Moreover, while the qualitative implications may be identical, this choice will have significant quantitative implications, which matter a lot in empirical work. So if you make this assumption you want to make sure that it matches your empirical setting. Note that [Horn and Wolinsky \(1988\)](#) is a theoretical paper. The first empirical application of their model was published much later ([Crawford and Yurukoglu, 2012](#)).
- The two options above are consistent with a setting where the supplier sends negotiating teams to the location of each downstream retailer and teams only find out about the outcome of other negotiations when they return to base. You can see how this is a bit artificial. Failure of a negotiation will have ripple effects on the other negotiation. Things that you may want to consider are:
 - Should the threat points change in the other negotiation?
 - Does the retailer have any recourse if negotiations fail? E.g. is there an “outside-option” supplier that allows them to keep operating, perhaps at a higher cost?
 - Does the supplier have any recourse? E.g. could there be a third retailer that could enter the market quickly if needed, and can we sell to them instead?

In general, accounting for these things is very complicated, because it requires building a lot of added structure. It’s also an area of active research at the frontier of the field. If you are interested, I cover some of those in the last section.

The most commonly used assumption is the first one, which lets us write the [Horn and Wolinsky \(1988\)](#) problem as

$$\max_{w_i} \left[\pi_i^*(w_i, w_j^S) \right]^{b_i} \times \left[w_i q_i^*(w_i, w_j^S) + w_j^S q_j^*(w_j^S, w_i) - w_j^S q_j^*(w_j^S, \infty) \right]^{1-b_i}$$

where the assumption on the threat point of the downstream retailer is that without an agreement they are out of the market.² After simplifying the FOC, one can get an implicit expression for w_i which can be combined with the analytical results from the downstream market as well as

²[Horn and Wolinsky \(1988\)](#) also assume symmetric bargaining weights (i.e. $b_1 = b_2 = \frac{1}{2}$), while the formulation above allows for asymmetric ($b_1 \neq b_2$) and heterogeneous ($b_i \neq \frac{1}{2}$) weights.

functional and distributional assumptions to get moment condition for estimation:

$$w_i = \frac{b_i w_j^S \frac{\partial \pi_i^*}{\partial w_i} \left(q_j^* \left(w_j^S, \infty \right) - q_j^* \left(w_j^S, w_i \right) \right) - (1 - b_i) \pi_i^* \left(w_i, w_j^S \right) \left(q_i^* \left(w_i, w_j^S \right) + w_j^S \frac{\partial q_i^*}{\partial w_i} \right)}{b_i \frac{\partial \pi_i^*}{\partial w_i} q_i^* \left(w_i, w_j^S \right) + (1 - b_i) \pi_i^* \left(w_i, w_j^S \right) \frac{\partial q_i^*}{\partial w_i}}$$

2.2 An example from health care: [Gowrisankaran, Nevo and Town \(2015\)](#)

We now consider an example from health care markets: the formation of hospital networks by insurance companies. Some background. In the US, health insurance plans have three basic components:

1. A monthly premium: this is what you pay every month regardless of utilization.
2. A cost-sharing schedule: this is what determines what fraction of overall expenditures are covered by the insurer, and what fraction has to be paid out of pocket.
3. A network of affiliated providers: this is a list of hospitals and physicians that are covered by the insurance company. Often, insurers also cover health care received out of network, though at much higher cost-sharing rates.

The latter element is where the structure of the market becomes vertical. Insurers have to negotiate rates with hospitals in order to include them in their network. The problem of how these networks are formed has come to the forefront of empirical IO research over the past 10-15 years, especially thanks to the work of Kate Ho (see e.g. [Ho, 2006, 2009](#); [Ho and Lee, 2017, 2019](#)).

Here we present a different paper, which was the first to look at bilateral bargaining in the context of hospitals and insurers ([Gowrisankaran, Nevo and Town, 2015](#)).

Setup

Their model has two stages. In the first stage, hospitals and insurers (which they call managed care organizations, or MCOs) bargain over prices. In the second stage, patients get sick and, with some probability, need to go to the hospital.

They make three assumptions that are quite limiting. The first is that networks are fixed.³ The second is that patients are assigned to insurers and cannot switch if they don't like the network. This works in their specific setting because they are looking at insurers that work closely with employers (that's what MCOs are). The assumption that employees will not switch employers because a hospital was dropped by the network is probably okay. The last assumption is that plan characteristics other than the network are fixed. This matters if you are worried that an insurer would want to set a higher copay for a hospital that is more expensive, as a way to reduce cost.

³As of yet, I don't believe there is a good model of hospital network formation that works broadly for empirical research. The paper that gets closest is [Ho and Lee \(2019\)](#), which is discussed in the last section. Even that paper however, does not estimate the network formation stage directly, opting instead for a simulation approach.

Their setup uses a fixed coinsurance, so the outcome of the negotiation does affect cost-sharing, though only in a mechanical way.

We skip the derivation of the first stage, which you can find in the paper. It's basically a hospital choice model where they recover some primitives that allow them to construct the expected utility of including an extra hospital in the network.

Bargaining between hospitals and insurers

They consider J hospitals, indexed by j , and grouped in $S \leq J$ systems, indexed by s , as well as M insurers, indexed by m . To build the objective function of the Nash bargaining problem they have to construct cost functions and profit functions. We present a simplified derivation here, which abstracts from heterogeneity in terms of disease. The paper has a more general setting that I encourage you to check out.

Each insurer sets a pair $(\mathcal{N}_m, \mathbf{p}_m)$ indicating a network $\mathcal{N}_m$, and a vector of negotiated prices with hospitals in the network $\mathbf{p}_m = \{p_{mj}\}_{j \in \mathcal{N}_m}$. For a given pair, the cost function of the insurer is

$$TC_m(\mathcal{N}_m, \mathbf{p}_m) = \sum_{i \in \mathcal{I}_m} (1 - c) f_i \sum_{j \in \mathcal{N}_m} p_{mj} s_{ij}(\mathcal{N}_m, \mathbf{p}_m)$$

where $\mathcal{I}_m$ is the set of patients that are insured by insurer m , c is the coinsurance rate, f_i is the probability that patient i will fall sick, and $s_{ij}(\cdot)$ is the probability a patient will choose hospital j . Given the cost, the net value (in dollars) of that same pair is

$$V_m(\mathcal{N}_m, \mathbf{p}_m) = \frac{1}{\alpha} \sum_{i \in \mathcal{I}_m} W_i(\mathcal{N}_m, \mathbf{p}_m) - TC_m(\mathcal{N}_m, \mathbf{p}_m) \quad (4)$$

where $W_i(\cdot)$ is the expected welfare of employee i from the pair $(\mathcal{N}_m, \mathbf{p}_m)$, and α is the price sensitivity parameter in the employee utility function (which allows us to transform utility units in dollars).⁴

Next, we need to build the profit function of each hospital system. A system has a set $\mathcal{M}_s$ of hospitals. They assume a constant marginal cost, but heterogeneous across hospitals and insurers:

$$mc_{mj} = \theta_j + \gamma_m + \varepsilon_{mj}$$

The justification here is that insurers may attract different patients, or the administrative requirements to submit claims for different insurers. In practice, keep in mind that these parameters is where the estimation hides unobserved heterogeneity.⁵ The number of patients that go to a

⁴The derivation in the paper allows for an additional parameter τ to account for the possibility that employer, employee, and insurer incentives are not perfectly aligned. We disregard that here.

⁵This is a bit harsh. "Hides" is a strong word with a negative connotation, which is not necessarily my intention. But you want to be aware of which parameters help absorb the differences between the model and the data, and be able to think critically about whether they are good representations of those differences. In this specific case, for example, could marginal cost heterogeneity hide something about unobserved variation in hospital profits by insurer (say, if

specific hospital is

$$q_{mj}(\mathcal{N}_m, \mathbf{p}_m) = \sum_{i \in \mathcal{I}_m} f_i s_{ij}(\mathcal{N}_m, \mathbf{p}_m)$$

Hence, the profits of a hospital system from a given set of contracts is

$$\pi_s(\mathcal{M}_s, \{\mathbf{p}_m\}_{m \in \mathcal{M}_s}, \{\mathcal{N}_m\}_{m \in \mathcal{M}_s}) = \sum_{m \in \mathcal{M}_s} \sum_{j \in \mathcal{J}_s} q_{mj}(\mathcal{N}_m, \mathbf{p}_m) \times [p_{mj} - mc_{mj}] \quad (5)$$

Finally, to specify the Nash bargaining problem, they make the (strong) assumption that contracts are unaffected by negotiating failures (this is the same as in [Horn and Wolinsky, 1988](#)). The problem is

$$\max_{\mathbf{p}_m} \left[\sum_{j \in \mathcal{J}_s} q_{mj}(\mathcal{N}_m, \mathbf{p}_m) \times [p_{mj} - mc_{mj}] \right]^{b_{s(m)}} \times [V_m(\mathcal{N}_m, \mathbf{p}_m) - V_m(\mathcal{N}_m \setminus \mathcal{J}_s, \mathbf{p}_m)]^{1-b_{s(m)}}$$

where $b_{s(m)}$ is the bargaining power of system s when facing insurer m , and $V_m(\mathcal{N}_m \setminus \mathcal{J}_s, \mathbf{p}_m)$ is the value of insurer m of a network that excludes system s (but includes all other systems in $\mathcal{N}_m$).

They solve the model by deriving the first-order conditions and use them to build moment conditions. The main FOC yields an equation of the form

$$p_{mj} = mc_{mj} - (\Omega + \Lambda)^{-1} q_{mj} \quad (6)$$

where the matrices Ω and Λ are functions of the model parameters (Ω is the first order derivative matrix for quantity with respect to price, and Λ is a function of bargaining weights and threat points).

Estimation

The parameters they need to estimate are:

- the bargaining weights $b_{s(m)}$
- the cost fixed effects γ_m and θ_j

They obtain a residual (the econometric error), by plugging the parametric form of the marginal cost in equation 6 to get

$$\hat{\varepsilon}_{mj}(b_{s(m)}, \gamma_m) = p_{mj} - \theta_j - \gamma_m + (\Omega + \Lambda)^{-1} q_{mj}$$

where I am writing the error term as a function because it depends on the chosen model parameters.

our price data was incomplete)? Or, what would happen if we misspecified functional form and marginal cost was increasing in quantity (because it's harder to run a hospital closer to capacity)?

The moment condition is

$$\mathbb{E} [\varepsilon_{mj} | Z_{mj}] = 0$$

where Z_{mj} is a vector of instruments. You need instruments here because Ω and Λ are functions of endogenous outcomes (prices and quantities), and therefore are correlated to ε_{mj} (lower ε_{mj} means low marginal cost, which in turn means lower equilibrium price and higher quantity). The estimation procedure involves iterating over guesses for $b_{s(m)}$, γ_m , and θ_j , and finding the combination that minimizes the moment condition.

Final notes

- Bargaining occurs over a single pricing dimension. This is obviously a simplification, since hospitals offer a bunch of services, at different prices. Using a linear price can be a good approximation if we think of the price as a multiplier for a full menu of fixed prices. But of course there can be exceptions. The paper provides some discussion of the approach and its limitations. If you are going to work in this area I encourage you to read the data section of the paper in detail.
- The identification of marginal cost and bargaining parameters is messy, because they have very similar implications for equilibrium: both higher bargaining weights and higher marginal costs imply higher prices, so how do we distinguish between the two? As far as I can tell, identification in this paper is largely by functional assumption. This is a limitation, but, to be fair, it would be very difficult to do better. Data on marginal costs is virtually impossible to identify, and bargaining weights are not really a data element. Ideally what you would want here is an exogenous cost shifter to use as instrument.
- Allan Collard-Wexler has another set of notes that goes over [Gowrisankaran et al. \(2015\)](#). You can find them [here](#).

3 RECENT ADVANCES IN THE EMPIRICAL ESTIMATION OF BARGAINING MODELS

As I mention in several places in the prior two sections, the big area of potential improvement in the derivation of bargaining models for empirical research is the creation of more realistic threat points. The difficulty in coming up with such an alternative can be boiled down to a series of requirements:

1. **Coming up with an “disagreement network”** that isn’t simply the current network, minus the hospital with whom the insurer is negotiating. This is actually not that hard. You could replace the hospital with a random one, or with the next “closest” hospital, according to some metric.

-
2. **Determining the price** for the disagreement network. This is harder. Applying Nash-in-Nash to the new network does not necessarily work, at least not in a straightforward way (what is the threat point for this new hypothetical equilibrium?). The solutions below have occasionally applied threshold rules, like TIOLI offers.
 3. **Checking that the new model has an equilibrium.** Because these are multiple-agent games, the existence (or uniqueness) of equilibria is not guaranteed.
 4. **Checking that estimation is feasible.** More complicated models are harder to estimate.

Below I list and briefly describe several research papers that have proposed innovations on the simple static Nash threat point. As a testament to the importance of health care markets in empirical IO, *all three* of these papers are about bargaining between insurers and hospitals.

Ho and Lee (2019)

This paper proposes what they call a Nash-in-Nash with threat of replacement (NNTR) solution. The bargaining stage is preceded by a network-formation stage where downstream retailers decide which suppliers to negotiate with. In the negotiation stage, the retailer can threaten to “replace” the supplier that they are negotiating with with someone from the list of suppliers that were left out of the network in the first stage. This solution has several appealing features. First, it provides a good explanation for why networks may be incomplete. Leaving suppliers out in the network-formation game leads to lower prices in the negotiation process. Notice that in standard Nash-in-Nash bargaining models partial networks are basically never optimal (the only exception is if there are really no gains from trade, which is generally unlikely), because excluding a supplier actually makes the other suppliers marginally more important, which in turn increases the equilibrium price to the detriment of the retailer. Second, it appears to be a more realistic description of what a negotiation might actually look like. The standard Nash framework models the retailer as a very passive agent: if negotiations fail there is no recourse. This is intuitively unrealistic, so it’s good to have a model that allows for some reaction.

There are some drawbacks however. The empirical application does not actually estimate the network formation game, because it’s too complicated (even in a setting with a single insurer and five hospitals). The price of the replacement threat is determined by a TIOLI offer from the insurer to the replacement hospital. This makes the replacement look pretty good (the insurer can capture all the surplus). Moreover, it is not clear that substitutability between the hospital that is threatened with replacement and its replacement matters all that much. Ghili (2020), which I cover below, argues that the framework allows for unrealistic threats, such as ones coming from hospitals in completely different markets.

Liebman (2018)

This paper uses a model that extends work by [Collard-Wexler et al. \(2019\)](#). [Collard-Wexler et al. \(2019\)](#) provides a micro-foundation to the Nash-in-Nash model of [Horn and Wolinsky \(1988\)](#) (using an approach that mimics the micro-foundation of the simple Nash bargaining problem in [Rubinstein, 1982](#)). They derive the result in [Horn and Wolinsky \(1988\)](#) as the equilibrium of a discrete, infinite-horizon dynamic game where upstream and downstream firms alternate offers within a pre-established network $\mathcal{G}$.

The additional element in [Liebman \(2018\)](#) is an initial stage where insurers (or, more generally, downstream retailers) choose the number of hospitals (or upstream suppliers) in their network. Then, in the alternating offer game, the structure of the network $\mathcal{G}$ can be stochastically altered by replacing a hospital that refuses an insurer offer with a randomly chosen hospital among the excluded ones. If an insurer excludes more hospitals the likelihood of replacement increases. [Liebman \(2018\)](#) argues that this structure allows for more realistic implications. He claims that in [Ho and Lee \(2019\)](#) most of the gains from exclusions arise from excluding a single hospital, whereas his model accommodates broader exclusion patterns that are sometimes observed. One could counter that random replacement seems unrealistic, though it may be possible to parametrize the replacement probability in a more realistic way (the paper does not do this as far as I can tell). This is a criticism that is pointed out also by the next paper, [Ghili \(2020\)](#).

Ghili (2020)

This paper proposes a third approach to extending Nash-in-Nash and moves several critiques to the two preceding papers. The most relevant critique is to the assumption in [Ho and Lee \(2019\)](#) that the insurer can commit to excluding a hospital even if exclusion is not incentive compatible (ex-post). The idea is that an insurer could exclude a hospital to threaten another, but then might still want to include the excluded hospital afterwards. The solution proposed by [Ghili \(2020\)](#) rules out these kind of exclusions. One possible issue with this approach (which is pointed out by [Ho and Lee, 2019](#) as well) is that it relies on unobserved fixed costs of “coverage” per hospital to explain why certain plans have narrow networks, while others don’t. There don’t seem to be many good practical reasons why it would be very costly for an insurer to cover many hospitals. Moreover, if the goal of the exercise is to show that exclusion is optimal for negotiating reasons, the presence of these fixed costs seems like a big confounder.

4 BIBLIOGRAPHY

- Collard-Wexler, Allan, Gautam Gowrisankaran, and Robin S. Lee**, "'Nash-in-Nash" Bargaining: A Microfoundation for Applied Work," *Journal of Political Economy*, 2019, 127 (1), 163–195.
- Crawford, Gregory S. and Ali Yurukoglu**, "The Welfare Effects of Bundling in Multichannel Television Markets," *American Economic Review*, apr 2012, 102 (2), 643–685.
- Dubois, Pierre, Ashvin Gandhi, and Shoshana Vasserman**, "Bargaining and International Reference Pricing in the Pharmaceutical Industry," 2019, (May), 1–75.
- Edgeworth, Francis Y.**, *Mathematical Psychics. An Essay on the Application of Mathematics to the Moral Sciences*. 1881.
- Feng, Josh**, "Pricing Intermediaries in Prescription Drug Markets: Impact and Policy Implications," 2019.
- Ghili, Soheil**, "Network Formation and Bargaining in Vertical Markets: The Case of Narrow Networks in Health Insurance," 2020.
- Gowrisankaran, Gautam, Aviv Nevo, and Robert Town**, "Mergers when prices are negotiated: Evidence from the hospital industry," *American Economic Review*, 2015, 105 (1), 172–203.
- Grennan, Matthew**, "Bargaining ability and price discrimination: Empirical evidence from medical devices," *American Economic Review*, 2013, 103 (1), 145–177.
- **and Ashley Swanson**, "Transparency and negotiated prices: The value of information in hospital-supplier bargaining," *Journal of Political Economy*, 2020, 128 (4), 1234–1268.
- Ho, Kate and Robin S. Lee**, "Insurer Competition in Health Care Markets," *Econometrica*, 2017, 85 (2), 379–417.
- **and —**, "Equilibrium provider networks: Bargaining and exclusion in health care markets," *American Economic Review*, 2019, 109 (2), 473–522.
- Ho, Katherine**, "The welfare effects of restricted hospital choice in the US medical care market," *Journal of Applied Econometrics*, nov 2006, 21 (7), 1039–1079.
- , "Insurer-provider networks in the medical care market," *American Economic Review*, 2009, 99 (1), 393–430.
- Horn, Henrick and Asher Wolinsky**, "Bilateral Monopolies and Incentives for Merger," *The RAND Journal of Economics*, 1988, 19 (3), 408–419.
- Liebman, Eli**, "Bargaining in Markets with Exclusion: An Analysis of Health Insurance Networks," 2018.

Nevo, Aviv, “Mergers that increase bargaining leverage,” 2014.

Rubinstein, Ariel, “Perfect Equilibrium in a Bargaining Model,” *Econometrica*, 1982, 50 (1), 97–109.