

ARMED CONFLICT BARGAINING

JAMES D. FEARON AND XIN JIN*

DRAFT

ABSTRACT. In armed conflict bargaining, states (or governments and rebel groups) may exchange offers while they are fighting, and conflict can end either with a negotiated deal or a military victory. If fighting is driven by private information about military capabilities or costs relative to benefits, why doesn't the possibility of frequent offers lead to rapid learning and convergence on a settlement, as tends to occur in standard buyer-seller bargaining models? We show that when private information concerns military capability, equilibrium war durations can be quite long for reasonable parameter values, and the conditional probability of negotiated settlement is U-shaped. Armed conflict bargaining with private information about capabilities is an interdependent values (aka "lemons") problem, but one that differs from buyer-seller contexts in that some fighting can directly reveal information about the value of more fighting. We also show that if, as makes sense for this setting, the sides can't commit to implement deals implying "ex post regret," then in equilibrium there can be long stretches of pure fighting with no exchange of serious offers. This is typically the case for both interstate and civil wars.

1. INTRODUCTION

At least since Clausewitz, war has been understood as a kind of bargaining process (Wagner, 2000). But what sort of bargaining process is it? One distinctive feature is that whereas in buyer-seller bargaining, the good in question changes hands only if both parties agree on the price, in *armed conflict bargaining* between two states or between a government and a rebel group, the "good" can change hands either by negotiated agreement or pure military victory. Indeed, combatants may understand the prospect of military defeat as the source of pressure that they hope will cause the other side to make concessions and negotiate a better deal. Wagner (2000) argues that this is in

Date: March 13, 2021.

* Fearon: Stanford University. Jin: ? We thank Avi Acharya, Gabriel Carroll, Massimo Morelli, Robert Powell, Guido Tabellini, and members of the Canadian Institute for Advanced Research IOG group for comments.

fact what Clausewitz meant by his famous statement that “War is the continuation of politics by other means.”¹

In this paper we clarify the nature of the widely employed analogy of war as a bargaining process, an idea referred to as “the bargaining model of war” in recent work (Reiter, 2003). We consider a model in which two parties – represented as a government and rebel group although they could be two states, or a union and firm – bargain while fighting. One side makes offers, rejection of which yields fighting that may result either in one side winning or in stalemate and continuation to another offer. The side receiving offers has private information about either its per period costs of fighting, or its military capabilities and prospects, or both. We consider what happens in equilibrium as we allow the offering side to make proposals in rapid succession, so as not to explain war by the time it takes to put offers together.²

When private information is only about fighting costs, in equilibrium the total war duration approaches zero as the time between offers grows small. Private information cannot explain significant delay (thus war) in this case, which proves to be analogous to buyer-seller bargaining with “independent values.” Here, the logic of the Coase conjecture obtains: Because the side making offers (say, the buyer) cannot commit not to quickly make a better offer after a rejection, it loses the ability to “screen” between types of seller, and trade (negotiated settlement) happens rapidly at prices close to the reservation value of the toughest type of seller (Coase, 1972; Gul, Sonnenschein and Wilson, 1986).

By contrast, if the side receiving offers has private information about its military prospects, then we find that empirically realistic war durations – even years – are possible in equilibrium for reasonable parameter values and as time between offers goes to zero. In this case, armed conflict bargaining is analogous to buyer-seller bargaining with *interdependent* values, meaning

¹Clausewitz distinguished between “absolute war” – a pure duel, or fight to the finish – and “real war.” In real war, the possibility or shadow of absolute war exerts pressure on the combatants to negotiate a political settlement. Hence war is a continuation of politics (seeking a political goal) by other means.

²Not to say this might not be relevant in some cases, but empirically it does not seem to be a dominant constraint on how long wars last. See discussion below, and Powell (2002).

that one party's private information affects the other's valuation for the good (Ausubel, Cramton and Deneckere, 2002). Bargaining over the price of a used car is the classic example (Akerlof, 1970).

Armed conflict bargaining is subtly different from the buyer-seller context in this case, however. In economic contexts, the seller has private information about the good being traded, which the buyer experiences only after trade is completed. By contrast, in armed conflict, one combatant's private information about the military situation affects the other combatant's value for *fighting*, which is the equivalent of "no trade" in the buyer-seller context. This has two interesting and consequential strategic implications. First, the act of fighting (delay, failure to trade) can itself directly reveal private information about the value of continued fighting, which in turn affects the parties' negotiating behavior in complex ways.

Second, buyer-seller models assume that if a price offer is accepted, the parties are committed and the deal is implemented, even though when there are interdependent values the buyer may suffer "ex post regret." That is, it can happen in equilibrium that acceptance of an offer reveals that the buyer has purchased a "lemon" not worth as much as she paid for it. Absent a warranty, the buyer has no recourse. By contrast, in the anarchical environments of armed conflict bargaining, the "buyer" *can always go back to war*. This means that deals that would imply ex post regret are not feasible. In Section 8 we modify the model so that the side making the offers has to "ratify" a deal accepted by the other side if it is to be implemented; alternatively it can go back to war. We find that equilibrium behavior in this *no-commitment case* changes in ways that match a set of puzzling stylized facts about armed conflict bargaining.³

In particular, the inability to implement deals that would be worse than war for one side proves to imply that armed conflict bargaining will be characterized by some degree of *fighting rather than bargaining*. In equilibrium the offering party may start out making serious offers that

³As we discuss below and as made clear by Deneckere and Liang (2006), ex post regret is also the reason that the Coase conjecture can fail (there is inefficient delay even in the limit) in buyer-seller bargaining with interdependent values. Coase conjecture dynamics operate when the buyer can't commit not to make another offer *and* has an incentive to accelerate trade because there is an expected surplus. But when offered prices rise to the level that imply ex post regret if accepted, the buyer's expected surplus from trade can be very small so she doesn't want to accelerate trade by making better offers. Long periods where there is very low probability of an accepted offer result.

have a positive probability of acceptance (and a stable peace deal), but quickly shifts to making non-serious offers that are commonly known to have zero chance of acceptance. Fighting proceeds until direct revelation of information about the military balance leads the offering side to suddenly shift to an offer that will certainly be accepted, so that a negotiated settlement eventually occurs if the conflict was not already decided by military victory. The implication is not only that the hazard rate of negotiated settlement in war will be U-shaped, which we show is true for the case of commitment to accepted offers, but also that the hazard rate of negotiated settlement can be zero for an extended period (“fighting rather than bargaining”), accompanied by non-serious offers that everyone knows have no chance of acceptance.

In the buyer-seller model, offers are always serious even if the chance of acceptance is small, and the offering side’s proposals increase monotonically. In armed conflict bargaining without commitment to offers, equilibrium may necessarily involve a (possibly long) period without any serious offers and zero chance of a negotiated settlement. It can also have a “price path” where the sides simply leave their offers, or war aims, unchanged while fighting goes on.

This pattern is a better match to what historians and political scientists have observed for both civil and interstate wars. They are essentially never characterized by constant back-and-forth and frequent revisions of proposals. Instead they often see long periods of fighting for completely incompatible wars aims – that are typically surprisingly unaffected by battlefield events – followed by sudden shifts to serious negotiation at a peace conference (Iklé, 1971; Reiter, 2009; Min, 2017). Our analysis suggests that this pattern may be a strategic consequence of combatants’ fear that making a concession will reveal weakness, leading to change of terms by the other side.⁴

The two closest articles in the literature are Deneckere and Liang (2006) and Powell (2004). Deneckere and Liang analyze buyer-seller bargaining with interdependent values. Our model differs by the incorporation of fighting, a different formulation of two type uncertainty, and the section

⁴Hart and Tirole (1988) identified and studied this “ratchet effect” problem for the context of buyer-seller bargaining with independent values and inability to commit against renegotiation (the example being bargaining over a rental good or service). There, the ratchet effect undermines screening and leads to immediate settlement (no delay). By contrast, with interdependent values and the ability to learn from fighting, it can yield long delays without serious bargaining.

that modifies the game to disallow commitment to deals that ex post would be worse for one side than return to fighting.

Powell studied essentially the same “bargaining while fighting” extensive form that we analyze.⁵ He allows for continuous types and considers the model with an arbitrary number of offers between discrete battles separated by a fixed amount of time. By contrast, in our version we allow fighting to be continuous while offers are made, or (an equivalent interpretation) the time between battles is allowed to be small. Just as here, Powell obtained a Coase conjecture result for the case where uncertainty is about costs of delay rather than military prospects; in his formulation no battles occur in the limit as the number of offers that can be made before a first battle increases. For the case of private information about military prospects, he shows there may have to be at least one battle, and that equilibrium takes the form of screening where both offers and battlefield results contribute to learning. However, he does not consider how much fighting the model can generate as offers can be made rapidly, or what the equilibrium path will look like in this case. In part because we analyze the simpler two-type case, we are able to characterize the equilibrium path and how the probabilities of military victory and negotiated settlement evolve over time, as well as demonstrating general conditions for failure of the Coase conjecture.⁶ Finally, like almost all “bargaining while fighting” models in the international relations literature, Powell assumed that states could commit to implement an accepted proposal, even if it is worse for the offering state than going back to war would be.⁷

Bester and Wärneryd (2006) and Fey and Ramsay (2011) apply the mechanism-design approach to armed conflict bargaining (Myerson and Satterthwaite, 1983). They both note that

⁵His extensive form has an additional move each period; the state receiving the offer can pass and put the onus of using force on the side making the offers. However, he also assumes that even the highest cost informed state is “dissatisfied,” meaning that it prefers fighting to the status quo, so this state never passes and the game is effectively like the one analyzed here.

⁶In buyer-seller models with or without interdependent values, updating with continuous types takes the form of truncation of prior distributions. When the private information is a probability of victory or survival from a period of fighting, the shape of a continuous distribution may change after every battle, which lowers tractability.

⁷See also Slantchev (2003) and Filson and Werner (2002). A precursor to this paper (Fearon, 2013), analyzed the no-commitment case but not (correctly) the commitment case. Fey, Meirowitz and Ramsay (2013) study bargaining in the shadow of a costly lottery (game-ending) war when offers can be rescinded.

private information about military capabilities implies interdependent values, and in turn the absence of an efficient mechanism that guarantees trade (zero probability of war).⁸ In both, war is characterized as a costly, one-time lottery, so they do not address war as a bargaining process or the question of how much fighting private information can produce. And again, as is necessary to apply the standard mechanism design approach, both assume ex ante commitment to agreements chosen by the mechanism, thus allowing commitment to deals that a state would want to renege on ex post. Our focus in this paper is on a bargaining game that can be analyzed under an assumption that rules out such deals; we comment briefly on novel issues involved in extending a mechanism design approach to this setting in Section 6.⁹

2. MODEL

A government G and rebel group R interact in periods $t = 0, 1, 2, \dots$. In each period G makes an offer x^t on the division of territory or revenues worth $\pi > 0$, a flow payoff. R then chooses whether to accept or reject. If R accepts the strategic interaction ends, with G receiving $\pi - x^t$ and R x^t in all subsequent periods. G and R have the same discount factor $\delta \in (0, 1)$ so that payoffs for an agreement on x^t are $(\pi - x^t, x^t)/(1 - \delta)$ from this period forward.

If R rejects, then the parties fight in this period (or continuous fighting continues to the next period) with three possible outcomes: Neither side wins militarily, which occurs with probability β ; R defeats G with probability α ; or G defeats R with probability γ , where $\alpha + \beta + \gamma = 1$. It will be convenient sometimes to use $\tilde{\alpha}$ and $\tilde{\gamma}$ as the exogenous parameters, where $\tilde{\alpha} = \alpha/(1 - \beta)$ and $\tilde{\gamma} = \gamma/(1 - \beta)$. $\tilde{\alpha}$ and $\tilde{\gamma}$ are the probabilities that R or G wins conditional on one of them winning in a period; note also that $\tilde{\alpha} + \tilde{\gamma} = 1$.

⁸See also Wittman (2009).

⁹Fey and Ramsay (2011) suggest an ex post participation constraint they call “voluntary agreements” into their mechanism design analysis; it appears not to play a role in their results, however. The mechanism design literature contains some work on trading problems with interdependent values that weaken the assumption of full ex ante commitment. For example, Forges (1999) considers “veto-incentive compatibility” (also Gerardi, Hörner and Maestri, 2014), and Shimer and Werning (2015) consider ex post participation constraints in the same manner as here, but focusing on issues of information transmission from the mechanism that do not arise in the two-type model that we analyze.

Regardless of the outcome, while they fight the two sides get reservation payoffs $g > 0$ and $r > 0$. We assume that $g + r < \pi$ so that fighting is costly relative to settlement. If one side wins, it takes control of π beginning in the next period while the loser is eliminated, so a victory by G gives it $g + \delta\pi/(1 - \delta)$ and likewise R gets $r + \delta\pi/(1 - \delta)$ if it wins. The loser's payoff is zero in all subsequent periods.

Two alternative interpretations of the model are worth noting. First, instead of a government and a rebel group fighting a civil war, it could depict an interstate conflict. Suppose state 1 makes offers to state 2 on the division of territory or other resources worth π . Let $q \in [0, \pi]$ be the status quo division, and let $g = \pi - q - c_1$ and $r = q - c_2$, where $c_i > 0$, $i = 1, 2$, are per-period costs of fighting. Or, the model could represent bargaining between a union and a firm over division of profits π , where we set $\alpha = 0$ and $\gamma = 1 - \beta$ is the per-period probability that the union leadership's ability to maintain the strike collapses.¹⁰

The complete information game has a unique subgame perfect equilibrium in which G always offers R 's per-period reservation value for always fighting, which is

$$\underline{x} = \frac{r(1 - \delta) + \delta\alpha\pi}{1 - \delta\beta},$$

and R always accepts any offer of at least this much.¹¹

Our concern will be with an incomplete-information version of the game in which R can be one of two types, "weak" and "strong." The prior probability that R is the weak type, R_w , is $\mu \in (0, 1)$ and G 's belief in period t is μ^t .

¹⁰Hart (1989) observed that Coase conjecture dynamics made it difficult to explain observed strike durations using plausible assumptions about discount rates. He obtained more realistic durations in a model that introduces a deadline effect (a known time when the firm's per-period probability of collapse increases markedly). Other strike models get realistic durations by assuming that one side can commit not to make or respond to offers for a given length of time, which rules out Coasian dynamics by fiat (Admati and Perry, 1987; Card, 1990). If strikes might arise not only due to uncertainty about firm profits but also (or instead) due to uncertainty about the union's (or perhaps firm's) ability to maintain internal discipline, then our results could also address the puzzle Hart posed.

¹¹The reservation value comes from $V_R = r + \delta\alpha(\pi/(1 - \delta)) + \delta\beta V_R$. Multiply by $1 - \delta$ to get the per period, or time-averaged values.

The types differ by their reservation values for fighting, with $\underline{x}_s > \underline{x}_w$. Substantively, $\underline{x}_s - \underline{x}_w$ is the gap that can make the government want to risk war by making tough demands $x^t < \underline{x}_s$ that a strong type would reject but a weak type might accept.

The difference between $\underline{x}_s$ and $\underline{x}_w$ can arise in multiple possible ways from differences in the two types' underlying parameters $M_i = (\alpha_i, \beta_i, \gamma_i)$ and r_i , where $i = w, s$ indexes the weak and strong types. The M_i parameters represent the *military capabilities* of the rebel group relative to the government, whereas r_i is the rebel group's flow payoff in a fighting period. $\underline{x}_s$ and $\underline{x}_w$ are given by the expression above with the r , α , and β terms subscripted by $i = w, s$.

Note that when $M_s \neq M_w$, the rebel group has private information about its military prospects relative to the government. This implies in turn that if the government is not sure whether it faces the weak or the strong type, it is unsure if its own reservation value for fighting is w_w or w_s , as given by

$$w_i = \frac{g(1 - \delta) + \delta\gamma_i\pi}{1 - \delta\beta_i}.$$

This is the case of *interdependent values*, since the rebel group's private information bears on the government's value for fighting.

By contrast, most of the bargaining literature on armed conflict considers models in which the parties' private information is not about military capabilities but instead costs for fighting (or value for the stakes), which are assumed to be *independent* across players. In our model, this means $M_w = M_s$ (there is agreement about military odds) but $r_s > r_w$ (the strong type has lower per period costs for fighting), which implies $\underline{x}_s > \underline{x}_w$. This case corresponds to buyer-seller models with independent values – for instance, the seller doesn't know the buyer's private value for the good.

In the case of interdependent values (private information about military capabilities), there are several ways that $M_s \neq M_w$ can give rise to $\underline{x}_s > \underline{x}_w$. To illustrate one empirically important case, in many small rural insurgencies there is essentially no chance that the rebels can overthrow the central government, but it is not clear if the government can decisively crush the rebels via successful

counterinsurgency.¹² Formally, $\alpha_s = \alpha_w = 0$ (rebels cannot take over the central government) and $\beta_w < \beta_s = 1$ (strong type cannot be defeated by counterinsurgency while there is a positive per-period chance the weak type can be). In this special case the two types' reservation values for war are thus $\underline{x}_w = r_w(1 - \delta)/(1 - \delta\beta_w)$ and $\underline{x}_s = r_s$.

In a second special case of interest, the stalemate probabilities are the same for both types, $\beta = \beta_w = \beta_s$, but the strong type has a greater per-period chance of decisive victory: $\alpha_s > \alpha_w$, implying also that $\gamma_s < \gamma_w$. Here the rebel groups' reservation values are

$$\underline{x}_w = \frac{r_w(1 - \delta) + \delta\alpha_w\pi}{1 - \delta\beta} \text{ and } \underline{x}_s = \frac{r_s(1 - \delta) + \delta\alpha_s\pi}{1 - \delta\beta}$$

To foreshadow, we find that in the case of private information about military prospects, much depends on the relationship between the stalemate probabilities of the two rebel types, β_s and β_w , because this determines what the government can learn *mechanically* from fighting. If $\beta_s = \beta_w$, then seeing the rebels survive a period of fighting tells the government nothing. In this case interdependent values can arise only if there is a difference between the strong and weak types' odds of winning decisively in a period (that is, $\alpha_s > \alpha_w$), and the government can learn about this difference only via different strategic behavior by the two rebel types in response to offers, not by observing battlefield outcomes. This case proves to be essentially equivalent to buyer-seller bargaining with interdependent values, or “lemon bargaining” (because bargaining over the used car reveals information about its quality only via the seller's behavior).

By contrast, if $\beta_s > \beta_w$, then seeing the rebel group survive a period of fighting mechanically increases the government's belief that it faces the strong type.¹³ This favors screening since, other things equal, rebel survival causes the government's belief to move in the direction that makes it more willing to make higher offers, that is, to “pool” weak and strong type because the strong type is more likely. If $\beta_s < \beta_w$, then the mechanical effect has the government *lowering* its belief that it

¹²Some of many possible examples include the current Afghan government and its NATO allies versus the Taliban in Afghanistan, Turkey versus the PKK, the Philippines versus Moro insurgent groups in Mindanao, or Myanmar versus a variety of regional insurgent groups.

¹³To push the “lemons” analogy, imagine that the seller and potential buyer bargain about the price of the used car *while driving it across the country* (and wear-and-tear effects are small).

faces the strong type as the rebel group survives longer, which increases the government's incentive to try to make lower offers. We will not analyze the $\beta_s < \beta_w$ case in this paper.¹⁴

Finally, we will often consider what happens as the time between offers, $\Delta > 0$, gets small. Let $\delta = e^{-\rho\Delta}$ and $\beta_i = e^{-\lambda_i\Delta}$, where $\rho > 0$ is a common discount rate and $1/\lambda_i$, $\lambda_i \geq 0$, is the expected duration of an “absolute war” (no settlements) between G and type R_i , $i = w, s$. As Δ goes to zero the probability of decisive victory in a very short span of time also goes to zero, so we have $\alpha_i = \tilde{\alpha}_i(1 - \beta_i) = \tilde{\alpha}_i(1 - e^{-\lambda_i\Delta})$ and likewise $\gamma_i = (1 - \tilde{\alpha}_i)(1 - e^{-\lambda_i\Delta}) = \tilde{\gamma}_i(1 - e^{-\lambda_i\Delta})$. So, when varying Δ , $\tilde{\alpha}_i = 1 - \tilde{\gamma}_i$, $i = w, s$, are the exogenous parameters describing military technology.

We should comment on what this formulation assumes about how fighting reveals information about relative capabilities. Powell (2004) fixed the time between battles and let the number of offers between battles get large, while assuming that there were military preparation or deployment costs paid while offers are exchanged. This is a reasonable approach for many interstate wars, one that pictures battles as discrete and sequential events that have a fixed minimum time separating them. However it does not allow us to ask what happens as the time between battles gets smaller, or what the implications are if fighting is conceived as a continuous activity that constantly reveals information (perhaps at a very slow rate). In the first article formulating bargaining while fighting in terms of a Rubinstein game with a risk of “break down,” Wagner (2000, 474) conceived of the military contest as “a continuous process taking place in the background while the exchange of offers takes some finite time to occur.” Many conflicts, both civil and interstate, see near-constant skirmishing, probing, bombing, or other military or guerrilla activities. So it can make substantive sense to see information revelation by the military channel as nearly continuous, even if very little might be revealed in a short amount of time. And even using Powell's interpretation, it seems worth knowing how many battles are needed to end a war, and how this number varies with observable parameters.

¹⁴The natural assumption for the guerrilla war example is the $\beta_s > \beta_w$, as discussed above. However, there is nothing particularly odd or impossible about $\beta_s < \beta_w$: For instance, suppose $\beta_s \ll \beta_w$ and α_s is much greater than α_w . This is a case where the strong rebel group is expected to be able to prevail quickly and take over the state, whereas a weak type can only manage a longer-running conflict with low odds of outright military victory.

3. EX POST REGRET AND G 'S COMMITMENT TO OFFERS

When the seller of a good has private information about its quality (like a used car), a buyer may suffer *ex post regret* if the seller accepts certain offers, since acceptance can reveal that the seller is trading a low-quality good not worth as much as what the buyer offered. The buyer would want to renege on the deal. The literature on bargaining with interdependent values typically rules reneging out under the assumption that the buyer can commit to pay *ex ante* via a contract.¹⁵

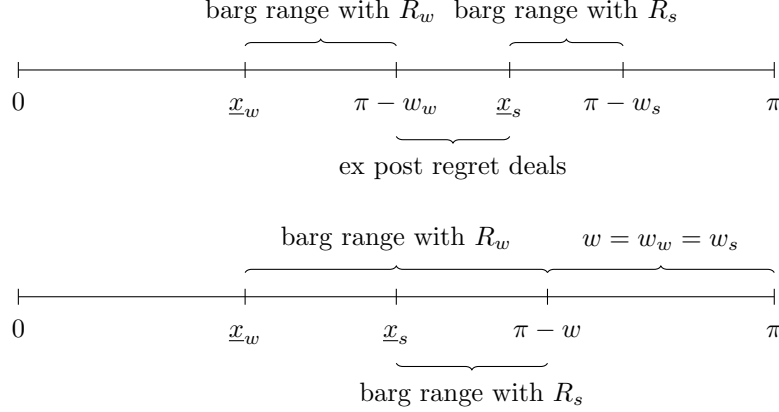
With private information about military capabilities the issue of *ex post regret* arises in armed conflict bargaining – the minimum offer needed to gain peace with a strong type can be more than the government would be willing to concede to a weak type because it sees its military prospects against a weak type as favorable. Thus there are deals that a strong type would reject, a weak type would in principle accept, but if accepted G would prefer to renege and go back to war.

This situation is shown graphically in the top line of Figure 1. Rebel group payoffs are measured to the right from zero and government payoffs to the left from π . To accept, the rebel group requires a deal $x \in [0, \pi]$ that gives it at least its (per period) reservation value for fighting, thus $x \geq \underline{x}_i$, $i = w, s$. If it knew which rebel type it faced, the government would require a deal x yielding a payoff $\pi - x$ at least as great as its reservation value for war w_i , $i = w, s$: thus $x < \pi - w_i$. So the bargaining range (the set of deals both prefer to war) would be $[\underline{x}_w, \pi - w_w]$ and $[\underline{x}_s, \pi - w_s]$ for weak and strong types, respectively. When war against a weak type is better for the government than war against a strong type ($w_w > w_s$), these ranges may not overlap, so that there are deals that if accepted would lead the government to want to renege and go back to war ($x \in (\pi - w_w, \underline{x}_s)$).

By contrast, in the independent values case shown on the second line in Figure 1, the government's value for fighting is the same against both types of rebel groups ($w_w = w_s$) so that any deal acceptable to the strong type of rebel group is also better for the government than continued fighting against the weak type. With independent values, any accepted deal that the government prefers to war is also *ex post* credible.

¹⁵But see Shimer and Werning (2015).

FIGURE 1. Bargaining ranges for interdependent (top) and independent (bottom) values cases



Virtually all buyer-seller bargaining models assume that the parties can commit to trade at an accepted price offer, which is natural given the assumed legal environment. But armed conflict bargaining occurs in the absence of third-party enforcement potential, and here a “buyer” with ex post regret can *unilaterally* decide to return to war. The government should be able, in effect, to reimpose “no trade” (war) at any time. As we will see, this will make screening by offers impossible for part of the bargaining process, to be replaced by “screening by fighting” instead.

4. EQUILIBRIUM IN THE GAME WITH COMMITMENT TO OFFERS

In this section we consider the game as described, which assumes that any offer accepted by the rebel group is implemented. We give a largely heuristic presentation of equilibrium and its properties, leaving full characterization and proofs to the Appendix.¹⁶

A history of play to period t is a sequence of offers $h^t = (x^0, x^1, x^2, \dots, x^{t-1})$. A complete strategy for G is a probability distribution on possible offers in $[0, \pi]$ for every t and every h^t (with $h^0 = \emptyset$), and a complete strategy for rebel group of type R_i , $i = s, w$, is a probability of rejecting

¹⁶So far as we know, a complete statement of equilibrium in the two-type, one-sided offer bargaining model has not been published. The difficulty lies in full characterization of off-path behavior and beliefs, which involve mixed strategies by both parties.

x^t , call it ν^t , for each t and each (h^t, x^t) . The equilibrium concept is perfect Bayesian with the additional requirement that strategies and beliefs satisfy “updating consistency” (Perea, 2002).¹⁷

Think of the government as the “buyer” who wants to buy peace from the rebel group in the form of a deal; x^t is the price offered to the rebel group, which is “selling” peace. Intuitively, unless G is very confident that it faces the tough type, G will start low and gradually raise the offer as R rejects and fighting proceeds without a victory (we are assuming $\beta_s \geq \beta_w$).

This is indeed what happens in the unique equilibrium. For a given set of parameters, bargaining/war will be finished in a maximum number of periods $N + 1$, thus period $t = N$. In the last period G offers $x^N = \underline{x}_s$ and the tough type accepts for sure. In the next to last period (if the war lasts this long), G makes an offer $x^{N-1} < \underline{x}_s$ that will be accepted for sure by the weak type and rejected by the strong type. Before this, the government makes an increasing sequence of offers x^t that are chosen so that R_w is indifferent between accepting x^t and rejecting, followed by accepting x^{t+1} if it survives to $t + 1$. It is straightforward to derive that this path is given by

$$(1) \quad x_n = (1 - (\delta\beta_w)^n)\underline{x}_w + (\delta\beta_w)^n\underline{x}_s,$$

where we adopt the convention that *time subscripts count down*, so that n here refers to the number of periods left until G makes the offer $x^N = \underline{x}_s$ and so ends the war for sure. Notice $x_0 = \underline{x}_s$, and $n = N - t$.

Until $t = N - 1$, the weak type mixes in response to offers on the path, rejecting with a probability $\nu^t \in (0, 1)$. (The strong type always rejects any offer less than $\underline{x}_s$.) As a result, the government’s belief that it faces the weak type declines along the equilibrium path, due both to strategy (the weak type is less likely to reject offers than the strong) and the mechanical effect (weak type is less likely to survive the fighting if $\beta_s > \beta_w$).

¹⁷Updating consistency is a slight refinement of PBE (weaker than sequential equilibrium) that imposes the minimum condition necessary for the one-deviation property to hold in an extensive form game (so that equilibrium can be checked by considering only single-period deviations). Notably it rules out PBE in which a deviation by, for instance, G in period t , leads G to have beliefs about R in $t + 1$ that are inconsistent with G ’s beliefs about R ’s type and behavior in period t . For example, in this game there can be PBE in which G makes non-serious offers on the path that both types of R reject, supported by G ’s expectation that if G deviated to a better offer, G would suddenly believe that R is certainly the strong type no matter what R did. Perea (2002) defined updating consistency for finite extensive forms; we extend it here to an infinite-horizon game.

An indifference condition pins down the equilibrium sequence of G 's beliefs and R_w 's rejection probabilities. In equilibrium the government must be indifferent between making the offer x_n when it has belief μ_n , and *skipping ahead* to the next, higher offer x_{n-1} . Skipping ahead has the advantage of saving on fighting costs and risks in order to “accelerate trade” (Deneckere and Liang, 2006), but the disadvantage of foregoing some probability that the weak type would accept the lower offer x_n . In equilibrium these costs and benefits balance at every time during the war.¹⁸

The sequence of equilibrium path beliefs μ_n is produced by a recursion beginning at the end. For example, μ_1 is the belief that the rebels are the weak type that is sufficiently low that G is indifferent between pooling both types with the good offer $\underline{x}_s$ (which both would accept), and making the lower offer x_1 that a weak type would be just willing to accept rather than wait a period to get $\underline{x}_s$. Then $\mu_2 > \mu_1$ is the belief such that G is indifferent between offering x_2 (which might be accepted or rejected), and skipping ahead to x_1 , which would induce separating. The weak type's equilibrium rejection probability ν_2 is exactly that which leads G to update from μ_2 to μ_1 using Bayes' rule.

We show that the equilibrium sequence of G 's beliefs (which imply as well the sequence of R_w 's rejection probabilities via Bayes' rule) is much more easily characterized using odds ratios. Let $o_n = \mu_n/(1 - \mu_n)$ and let the initial odds that G thinks R is the weak type be $o = \mu/(1 - \mu)$. Let $A = (\pi - w_w - \underline{x}_w)/(\underline{x}_s - \underline{x}_w)$, a constant that is the ratio of surplus when $R = R_w$ to the difference between the strong and weak types' reservation values for fighting. Define the sequence o_n , $n = 0, 1, \dots$, as follows: $o_0 = 0$, $o_1 = A(\pi - r_s - g)/(\pi - r_w - g)$, and for $n > 1$

$$(2) \quad o_n = o_{n-1} \frac{\beta_s}{\beta_w} (1 + A_n) - o_{n-2} \frac{\beta_s}{\beta_w} A_n,$$

¹⁸Why, along the equilibrium path, must G 's posteriors μ^t land at the cutoff values μ_n ? In equilibrium G has to offer the x_n sequence, or else R_w would not be indifferent between accept and reject. If G deviates to $x^t \neq x_n$, then the equilibrium of the continuation must involve R_w being indifferent between accepting and not accepting today (otherwise, R_w either accepts today with probability 1 or rejects today with probability 1, but neither of these can support continuation equilibria). In turn, R_w is indifferent only if G mixes in the next period following R 's rejection of x^t . But for G to mix in $t + 1$, G 's posterior μ^{t+1} must be at a cutoff value. So if G deviates to $x^t \neq x_n$, $\mu^{t+1} = \mu_k$ for some cutoff belief. This means that on the path, where G offers x_n , R_w still has to accept with a probability that makes G 's posterior tomorrow equal the cutoff; otherwise G would strictly gain today by offering slightly more or less.

where $A_n = A/(\delta\beta_w)^{n-1}$.

Since $A > 0$, it is easy to show that when $\beta_s \geq \beta_w$, the sequence o_n increases strictly and without bound, so that for any initial odds o there is a smallest N such that $o_{N+1} \geq o$. That N will be the maximum possible duration of the conflict.

The size of the term A strongly influences the rate at which o_n increases and thus the duration of armed conflict. Note that $A \geq 1$ iff $\pi - \underline{x}_s \geq w_w$, which means that G has no ex post regret for any deal $x^t \leq \underline{x}_s$; this condition is implied by independent values. And $A < 1$ only with interdependent values.

If $A \geq 1$ then for $n > 1$, o_n is increasing at an increasing rate. We show in the proof of Proposition 2 below that in this case the Coase conjecture holds and maximum duration approaches zero as the time between rounds gets small. By contrast, if $A < 1$, then o_n can increase at a decreasing rate at first, and may level out, increasing slowly for a long period of time before explosive growth finally kicks in. For large enough initial μ , the Coase conjecture fails.

We can now give a statement of equilibrium path strategies and beliefs in proposition form.¹⁹ Let $\mu_0, \mu_1, \dots, \mu_n$ be the sequence of government beliefs implied by o_n .

Proposition 1. *The game has a perfect Bayesian equilibrium that satisfies updating consistency. The equilibrium path of any such equilibrium is unique, except for G 's choice in the first period in a non-generic case mentioned below. This path is described as follows. For initial belief $\mu \in (0, 1)$ that the rebel group is the weak type, let N be the largest integer such that $\mu_N \leq \mu$.*

- *On the equilibrium path, G offers $x^t = x_{N-t}$, according to (1), for periods $t = 0, 1, \dots, N$.*
- *R_s rejects any offer less than $\underline{x}_s$ and accepts otherwise (both on and off the path).*
- *On the path R_w accepts for sure $x_1 = x^{N-1}$ in period $t = N - 1$. For $0 < t < N - 1$, R_w rejects offer x^t with probability $\nu^t = \nu_n \equiv \frac{o_{n-1}\beta_s}{o_n\beta_w}$, where $n = N - t$, and $\nu^0 = \nu_N = \frac{o_{N-1}\beta_s}{o\beta_w}$.*

¹⁹Off-path strategies and beliefs are complex, because deviations by the government to a larger or smaller offer can require that the government mix in multiple subsequent periods. Imagine, for instance, that G deviates to a slightly higher x^t . If R_w were to accept for sure, G would offer $\underline{x}_s$ in the next period, implying that R_w would not want to accept x^t for sure. R_w must mix in response to the better offer off the path, anticipating that G will next mix between the old x^{t+1} and the skipping ahead offer x^{t+2} so that R_w is indifferent. We provide details in the Appendix.

- G 's beliefs evolve along the path by $\mu^t = \mu_{N-t}$ for $t > 0$.

This equilibrium path is unique except in the non-generic case of $\mu = \mu_N$, where G is indifferent between offering x_N and x_{N-1} and so can play according to any mix between them.

5. WAR DURATION IN ARMED CONFLICT BARGAINING WITH COMMITMENT

What are the qualitative features of armed conflict bargaining as it occurs in equilibrium in this “commitment case”?

Consider first an independent values example. The pie is normalized to size 1, and we use $r_s = .5$, $r_w = .1$, and $g = .3$. This means that per-period costs of conflict are 20% of the size of the pie in a war with a strong type, and 60% in a war with a weak type. The annual discount rate is 5%, and we set the initial probability that the rebel group is the weak type to $\mu = .75$. The offer rate is set to one per month. Alternatively one can interpret the offer rate as the rate of battles, if one assumes one offer after each battle.

Suppose that $\lambda_w = \lambda_s = 1$ per year, which means that both types survive a month of fighting (or a battle) with probability $\beta_i = e^{-1/12} = .92$. Let $\tilde{\alpha}_w = \tilde{\alpha}_s = .5$, so that if the conflict ends with a military victory, G and R are equally likely to prevail. These assumptions imply independent values since the government's payoff for fighting is the same against both types, $w_i = .49$. The difference between the two types' reservation values for war proves to be small on a per-month basis, $\underline{x}_s = .5$ and $\underline{x}_w = .48$. This stems from the combination of discounting and the fact that the benefits of winning are long-run relative to the costs of fighting. So, in this case $A = (1 - .49 - .48)/.02 = 1.5$.

Table 1 shows median, mean, and maximum possible war durations and some other outcomes for this case (panel A row 2, or “A2”) and a range of others.²⁰ We see that if the initial offer is rejected the mean duration of war is less than two months ($.16 * 12 = 1.9$) and the conflict lasts at most three months. If we increase the offer (or battle) rate to one per day, the maximum duration is three days. Because fighting time is so short relative to the odds that a battle is decisive – note that $\lambda_i = 1$ means that a Clausewitzian “absolute war” would have mean duration of one year – the

²⁰Parameter values are the same except as varied in columns 1-5 or noted. Durations are conditional on war occurring.

conflict almost certainly ends by negotiated settlement. This case illustrates the Coase conjecture at work in a case of independent values, a result we prove more generally below and that parallels the standard result for buyer-seller bargaining with independent values.²¹

Next consider what happens if we introduce a small amount of private information about R 's military capabilities, beginning with a case of no direct learning from fighting ($\beta_w = \beta_s$). Let $\tilde{\alpha}_s = .51$ and $\tilde{\alpha}_w = .49$, so that we make the strong type slightly more likely to prevail if there is a military victory, and the weak type slightly less. Now the expected duration of a war more than doubles from two to 4.4 ($= .37 * 12$) months, and maximum duration triples from three to nine months, which is realized with an ex ante probability of approximately $.25 * .92^9 = .12$, the probability of the strong type times the probability that no military victory occurs in nine months.

Why the dramatic change? Mathematically, the reason is that even this small divergence in rebel types' military prospects leads to a large fall in A , from 1.5 to .77. The bargaining range with R_w has shifted down slightly – G has gotten stronger relative to R_w – while the bargaining range with R_s , including $\underline{x}_s$, shifts slightly higher, so increasing G 's possible gains from separating types, $\underline{x}_s - \underline{x}_w$. Critically, A is now less than one. This makes for a substantive change in equilibrium behavior related to the fact that there is now an ex post regret zone, $(\pi - w_w, \underline{x}_s)$, where an accepted offer makes G worse off than would more war.

Why does G even make offers in this range, if acceptance would be regretted? *Allowing some chance of a loss if R is the weak type is the only good way G can get to a surplus in a deal if R is the strong type.* Ideally, G might want to stop increasing offers at the time $t = N - \hat{n}$, where $x_{\hat{n}} \approx \pi - w_w$, so attempting to commit to “stand pat” long enough that the weak type would prefer to accept in this round. But the commitment is not credible since if R_w accepts $x_{\hat{n}}$, then G knows that rejection implies R_s , and so would immediately offer $\underline{x}_s$. R_w would want to mimic R_s to get this. G 's remaining option would be to make the pooling offer ($\underline{x}_s$) right away (at $t = N - \hat{n}$), but at this point μ^t is still high enough that G prefers to fight on.

²¹The case in Table 1, A1 is in fact exactly analogous to a buyer-seller model, since here the military technology is such that neither side can gain the prize by defeating the other militarily.

TABLE 1. War durations (in years) for different cases

A. Independent values ($r_s > r_w, \beta_s = \beta_w, \tilde{\alpha}_s = \tilde{\alpha}_w$)											
offers/year	$1/\lambda_w$	$1/\lambda_s$	$\tilde{\alpha}_s$	$\tilde{\alpha}_w$	A	duration war			P(war)	P($\cdot$ war)	
						median	mean	max		deal	mil vict
12	∞	∞	$\cdot$	$\cdot$	1.5	0.17	0.18	0.25	0.56	1	0
12	1	1	0.5	0.5	1.5	0.17	0.16	0.25	0.58	0.84	0.16
52	1	1	0.5	0.5	1.5	0.04	0.04	0.06	0.57	0.96	0.04
365	1	1	0.5	0.5	1.5	0.01	0.01	0.01	0.56	0.99	0.01
B. Interdependent values with no learning while fighting ($\lambda_w = \lambda_s \Rightarrow \beta_w = \beta_s$)											
offers/year	$1/\lambda_w$	$1/\lambda_s$	$\tilde{\alpha}_s$	$\tilde{\alpha}_w$	A	duration war			P(war)	P($\cdot$ war)	
						median	mean	max		deal	mil vict
12	1	1	0.51	0.49	0.77	0.33	0.37	0.75	0.78	0.65	0.35
365	1	1	0.51	0.49	0.75	0.23	0.27	0.55	0.81	0.73	0.27
12	1	1	0.75	0.25	0.06	0.75	1.01	5.17	0.26	0.03	0.97
12	1	1	1	0	0.03	0.75	1.02	5.83	0.26	0.02	0.98
12	1	1	1	0	0.01	0.75	1.02	7.67	0.26	0.02	0.98
C. Interdependent values with learning while fighting ($\lambda_w > \lambda_s \Rightarrow \beta_w > \beta_s$)											
offers/year	$1/\lambda_w$	$1/\lambda_s$	$\tilde{\alpha}_s$	$\tilde{\alpha}_w$	A	duration war			P(war)	P($\cdot$ war)	
						median	mean	max		deal	mil vict
12	1	∞	0	0	0.06	0.5	1.23	2.92	0.77	0.73	0.27
52	1	∞	0	0	0.06	0.29	1.02	2.73	0.85	0.79	0.21
12	1	∞	0	0	0.06	2.58	1.75	2.58	0.85	0.77	0.23
12	3	∞	0	0	0.03	1.33	3.42	8.67	0.87	0.69	0.31
52	3	∞	0	0	0.03	0.94	3.02	8.54	0.97	0.73	0.27

Notes: In all cases except as noted next, $\pi = 1, g = .3, r_s = .5, r_w = .1, \mu = .75, \rho = .05$ (per year). B, row 5: $\rho = .01$. C, row 3: $\mu = .5$. C, rows 4-5: $g = r_s = r_w = .45$.

To screen, G has to make offers in the ex post regret range, and they need to follow a particular path (equation 1) if R_w is to be willing to mix.²² The question of expected conflict duration concerns R_w 's probability of accepting offers in or near this range. Intuitively, acceptance must be unlikely because the deal would be as bad or worse for G than just fighting; if acceptance were likely, G would want to deviate to worse offers (i.e., ones not increasing as fast). By contrast, with independent values, R_w 's acceptance is not so bad because G realizes a surplus in any deal that R_w (or R_s) is willing to accept. In that case, G 's incentive is to "accelerate trade" (Deneckere and

²²Clearly equilibrium cannot involve R_w accepting with probability one before the next to last round. We show that in the proof that neither can R_w can reject with probability one in an equilibrium.

Liang, 2006): The surplus with both types combines with G 's inability to commit not to quickly make a better offer to yield the Coase conjecture.

For interdependent values ($A < 1$), the approximate maximum number of bargaining/war periods from the first offer in the ex post regret region, $x_{\hat{n}}$ to $x_0 = \underline{x}_s$ can be found from

$$\begin{aligned}\pi - w_w &\approx x_{\hat{n}} = (1 - (\delta\beta_w)^{\hat{n}})\underline{x}_w + (\delta\beta_w)^{\hat{n}}\underline{x}_s \\ (\delta\beta_w)^{\hat{n}} &\approx \frac{\pi - w_w - \underline{x}_w}{\underline{x}_s - \underline{x}_w} = A \\ \hat{n} &\approx \frac{\log 1/A}{\Delta(\rho + \lambda_w)}\end{aligned}$$

The length of time is thus

$$\hat{T} \equiv \hat{n}\Delta = \frac{\log 1/A}{\rho + \lambda_w}.$$

Total maximum war duration will therefore be less than $\hat{T}$ if $N < \hat{n}$, and greater if $N > \hat{n}$. Deneckere and Liang showed that in the two-type, buyer-seller interdependent values case (analogous to $\beta_w = \beta_s$ here) with large enough initial belief μ , the limiting maximum duration is $2\hat{T}$. $N > \hat{n}$, and it takes just as long to get from x_N to $x_{\hat{n}}$ as from $x_{\hat{n}}$ to x_0 .²³ The reason is that in the buyer-seller interdependent values case, equilibrium beliefs evolve symmetrically on either side of $\hat{n}$, so doubling the total time in the limit.

We find that this is no longer true with direct learning from fighting, or $\beta_s > \beta_w$. Beliefs don't evolve symmetrically around $\hat{n}$. We show that for initial beliefs greater than a threshold value that depends on parameters, limiting maximum war duration is greater than zero but at most $\hat{T}$, exactly half of the maximum duration in the analogous buyer-seller case. As discussed below, mechanical learning from fighting and strategic learning complement each other, making for more rapid change in G 's beliefs as R survives more war.

²³In their notation, maximum duration is $T = 2\log(\bar{c}/\underline{v})/r$, where r is the discount rate; $\bar{c}$ is the difference between the high and low quality sellers' costs of production of the good; and $\underline{v}$ is the surplus when the seller is the low quality type. In the armed conflict setting, the weak type's per period hazard of military defeat adds to the denominator, and military prospects are also incorporated in the terms in the numerator because these depend on payoffs from fighting. We provide a full proof of the $2\hat{T}$ result for our game in the Appendix.

Figure 2 illustrates these varied dynamics. Parameters are same as above (e.g., 12 offers per year), except as noted in the titles. Panel A is the independent values cases already discussed, in which war has maximum duration of three months. Panel B is an interdependent values case with no direct learning from fighting. Notice that R_w accepts the initial offer, which would avoid war, with a relatively high probability of about .9, and the second offer as well (about .8). Thus two rejections and a little fighting lead to a big change in G 's beliefs, after which R_w becomes extremely unlikely to accept an offer. G 's beliefs then change very slowly for a symmetric length of time, $\hat{n}\Delta$, on either side of the time where G offers $\pi - w_w$, when offers move into the ex post regret zone. In this example it takes about 4.5 years for G to get to the point where it is willing to make the pooling offer to end the war. Since the hazard rate on military victory is 1 for both rebel types, this means that this conflict would very likely end on the battlefield rather than at the table.

Though not shown explicitly, one can guess from the path of rejection/acceptance probabilities that the hazard of conflict termination in this example is U-shaped: Negotiated settlement is most likely either in the first few days or months, or at the maximum duration time, after a spell of fighting during which the chance of an accepted offer is very low. This proves to be a general feature of armed conflict bargaining in this two-type model, whether we are looking at a case of independent values, or the two varieties of interdependent values (learning from fighting, and no learning from fighting).

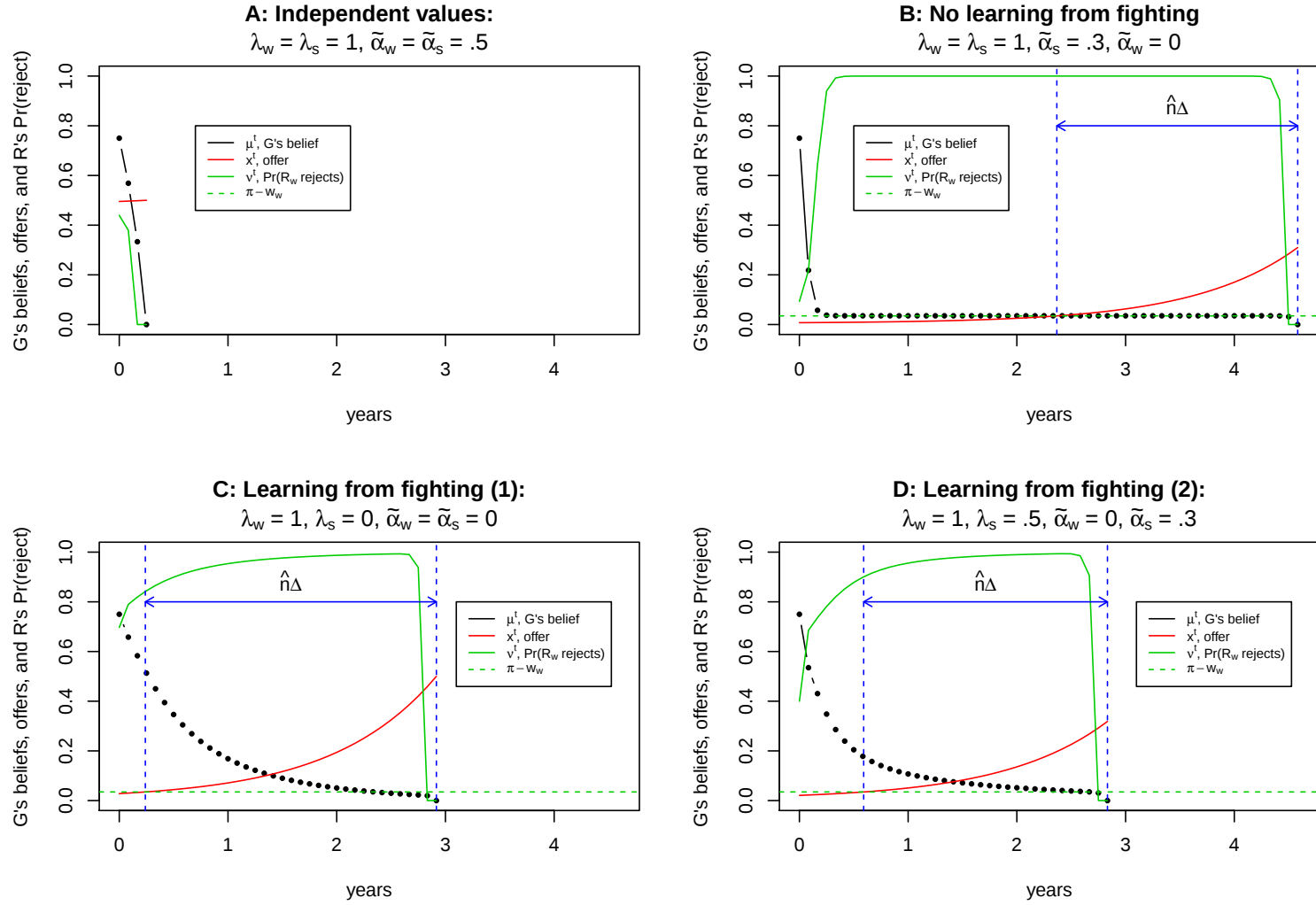
Panel C in Figure 2 illustrates an insurgency, in which the rebels have private information that bears on whether they can survive the government's counterinsurgency effort. Panel D is similar but gives the strong rebel group a .3 chance of winning, while increasing the hazard of decisive victory to $\lambda_s = .5$ per year. In both cases we see that the maximum duration falls towards $\hat{T}$, with the time it takes to get from N to $\hat{n}$ shrinking relative to panel B. Note that again the hazard rate of negotiated settlement is U-shaped.

Further, there is an interesting contrast with the "lemons bargaining" case (panel B) in terms of the evolution of the weak rebel type's probability of rejecting offers. When $\beta_w = \beta_s$, R_w is quite likely to accept the first couple of offers, and then quickly shifts to rejecting almost for sure. When

$\beta_s > \beta_w$, by contrast, R_w is more likely to reject the first few offers – so war is more likely to begin – but then only gradually increases its rejection probability towards 1, so making for a shorter overall duration (other things equal). We show in the next section that this is general for the model with small enough time between offers. The intuition is not completely clear, although it may result from G getting additional bargaining leverage from the fact that war itself can “screen” types of R .

Lastly, we note that for what seem like reasonable parameter values, examples in Figure 2C-D and Table 1C yield mean and maximum war durations that are close to empirical means and maximums. This model is too stylized and sparse to imagine calibration as meaningful, but it is nonetheless interesting that interdependent values can generate such long delay in armed conflict bargaining in principle. Other main features also match stylized facts about armed conflict bargaining. For example, hazard rates of conflict termination by negotiated settlement are U-shaped, and there are extended periods of fighting with very little chance of negotiated agreement till the conflict becomes “ripe for resolution” (e.g., Iklé, 1971; Zartman, 1989; Reiter, 2009).

FIGURE 2. Sample equilibrium paths



6. THE COASE CONJECTURE, LIMITING WAR DURATIONS, AND INEFFICIENCY

This section provides a proposition characterizing limiting maximum war duration as time between battles or offers becomes small for our several cases.²⁴ We also consider insights from a continuous approximation.²⁵

Proposition 2. *Let Δ be the time between offers in the armed conflict bargaining game when $\lambda_s \leq \lambda_w$, and let $N(\Delta)$ be the maximum number of rounds of bargaining in equilibrium. Then the total possible maximum duration of war, $N(\Delta)\Delta$, approaches zero as Δ approaches zero if and only if either*

(a) $\pi - \underline{x}_s \geq w_w$, which is implied by $A \geq 1$, or

(b) $\pi - \underline{x}_s < w_w$ and $\mu < \mu^*$ defined by

$$o^* = \frac{\mu^*}{1 - \mu^*} = \frac{o_1}{1 - A} = \frac{\rho + \lambda_s}{\rho + \lambda_w} \frac{\pi - w_s - \underline{x}_s}{w_w - (\pi - \underline{x}_s)}.$$

If (c) $\pi - \underline{x}_s < w_w$ and $\mu > \mu^*$, then the maximum and mean war durations are bounded away from zero in the limit as $\Delta \rightarrow 0$. When $\lambda_s = \lambda_w$ (i.e., $\beta_s = \beta_w$), limiting maximum duration is $2\hat{T}$, where $\hat{T} = \log(1/A)/(\rho + \lambda_w)$, for any $\mu > \mu^*$. When $\lambda_s > \lambda_w$, limiting maximum duration approaches $\hat{T}$ from above, but decreases as the initial belief μ falls towards μ^* .

We conjecture that in case (c) when $\lambda_w > \lambda_s$ (i.e., $\beta_s > \beta_w$), limiting maximum duration is given by the $T^* > 0$ that solves, for initial belief $o > o^*$,

$$(3) \quad o = o^* e^{(\lambda_w - \lambda_s)T} \left[\frac{1 - A}{1 - A e^{(\lambda_w + \rho)T}} \right]^{\frac{\lambda_w - \lambda_s}{\lambda_w + \rho}}.$$

See discussion below.

In words, the Coase conjecture applies, and mean or maximum war duration go to zero with the time between offers, if uncertainty concerns only costs of fighting (independent values) or if the initial belief that the rebel group is the weak type is sufficiently low ($o < o^* = o_1/(1 - A)$).

²⁴We do not mean to suggest that war duration as the time between offers is allowed to get arbitrarily small is the most empirically relevant question to ask about properties of equilibrium. But it is interesting from a theoretical perspective and very useful for insight into the strategic dynamics.

²⁵In this proposition and as relevant in those that follow, terms that depend on Δ , the time between offers, refer to their limiting values. For example, $\underline{x}_i$ refers to $\lim_{\Delta \downarrow 0} \underline{x}_i$, and likewise for w_i , o_1 , and A .

If the government's initial belief is above this threshold and the rebel group's private information concerns military prospects, then there will be a war with positive probability. In the buyer-seller interdependent values case ($\beta_w = \beta_w$), its maximum duration is $2\hat{T}$, which is independent of the initial belief and is increasing in the size of the ex post regret zone. In the "learning from fighting" case, the limiting maximum duration is at most $\hat{T}$, and decreases as the government's initial belief decreases that it is facing the weak type of rebel group.²⁶

Some insight into strategic differences between the buyer-seller interdependent values case and "learning while fighting" can be obtained by examining a continuous-time approximation of o_n , call it $o(\tau)$, where $\tau \geq 0$ is continuous time measured backwards from the maximum duration ($n = 0$). Rewriting the main recursion and using $r = \beta_s/\beta_w$, we have

$$o_n - o_{n-1} = rA_n(o_{n-1} - o_{n-2}) + (r-1)o_{n-1}.$$

For small Δ , if o_n is not increasing at a rapidly increasing rate, then the approximation $o_n - o_{n-1} \approx o_{n-1} - o_{n-2}$ will not be bad. So consider the alternative sequence

$$o_n - o_{n-1} = rA_n(o_n - o_{n-1}) + (r-1)o_{n-1}.$$

For brevity, let $L \equiv \lambda_w - \lambda_s$ and let $\rho' \equiv \rho + \lambda_w$. Dividing by the time interval Δ and taking limits as $\Delta \rightarrow 0$, we get

$$\begin{aligned} o'(\tau) &= \frac{L}{1 - Ae^{\rho'\tau}} o(\tau) \\ \frac{d}{d\tau} \log o(\tau) &= \frac{L}{1 - Ae^{\rho'\tau}} \\ o(\tau) &= o^* e^{L\tau} \left[\frac{1 - A}{1 - Ae^{\rho'\tau}} \right]^{\frac{L}{\rho'}} \end{aligned}$$

The last step is from integration and algebra using the initial condition $o(0) = o^*$, which is shown in the proof of Proposition 2.²⁷

²⁶Since λ_w and λ_s figure into A , comparative statics on these military parameters can be subtle. It is not difficult to show, however, that in the "pure insurgency" case where $\hat{\alpha}_i = 0$, increasing the gap $\lambda_w - \lambda_s$ reduces maximum duration.

²⁷The first step, from the discrete differences to $o'(\tau)$, also requires justification. See discussion in the Appendix. Because $o_n - o_{n-1} > o_{n-1} - o_{n-2}$ after a period of time that goes to zero as $\Delta \rightarrow 0$, $o(\tau)$ increases faster than $o_{\lfloor \tau/\Delta \rfloor}$. However, it appears that for any $\epsilon \in (0, \hat{T})$, the approximation error appears to go to zero with Δ for all $\tau \in (0, \hat{T} - \epsilon)$. We have so far not been able to prove this.

This analysis works only where $o'(\tau)$ is well defined, which is not the case when $\lambda_w = \lambda_s$ at times $\tau = 0$ and $\tau = 2\hat{T}$. There, $\lim_{\Delta \downarrow 0} o_n$ approaches a step function that rises at $\tau = 0$ to o^* , and is then flat until $2\hat{T}$ when it again goes vertical. Looking from the start of a possible war, this implies that a deal is struck at approximately $t = 0$ without any real conflict or delay with probability $\mu(1 - o^*/o)$. If conflict does start, the weak type of rebel group is extremely unlikely to accept an offer till war duration $2\hat{T}$ has passed, if the war goes on that long. This is so even though G is continually improving its offer. Then all of a sudden the probability of accepting offers suddenly increases and the conflict ends.

If there is learning from fighting ($\lambda_w > \lambda_s$), then in the limit o_n behaves the same way around $\tau \approx 0$ (rising immediately to o^*), but then increases at an increasing rate, tending towards vertical as it approaches $\hat{T}$. The increase in $o(\tau)$ – or, from the perspective of the start of the conflict, the decrease in G 's belief that it faces R_w as a war proceeds – is due to both the mechanical effect of learning while fighting *and* the strategic effect coming from R_w 's possible acceptance of an offer.

The mechanical effect is captured by the term $e^{(\lambda_w - \lambda_s)\tau}$ in $o(\tau)$, while the strategic effect is seen in the term in brackets. The rate of change of G 's log posterior odds that it faces the weak type – $\frac{d}{d\tau} \log o(\tau)$ above – can be written as

$$L + L \frac{Ae^{\rho'\tau}}{1 - Ae^{\rho'\tau}},$$

which shows that the elasticity of change in G 's beliefs is the sum of the mechanical component ($L = \lambda_w - \lambda_s$), and the strategic component, which reflects the rate at which R_w accepts offers along the path. The latter term has two interesting implications. First, more mechanical learning (greater L) implies that R_w also accepts at a higher rate, so that strategic learning also happens faster. Second, the rate at which R_w accepts is highest at the start of the bargaining process, and decreases until just before the end ($\tau = 0$). Thus, the hazard rate of negotiated settlement is decreasing, until it shoots up just before the final round.

The hazard rate of negotiated settlement is also U-shaped in the $\lambda_w = \lambda_s$ case, but much more extreme: a significant chance of a peace deal agreement occurs right away with scarcely any

fighting, or after a long period with almost no chance of an agreement. By contrast, learning while fighting makes military conflicts more likely to begin but also more likely to be settled by a deal if the fighting gets going. An intuition for the difference: $\beta_s > \beta_w$ reduces the weak type's ability to mimic the strong type, and thus its incentive to maintain a reputation for possibly being the strong type. This favors the government, in effect giving it some additional structural bargaining power.

7. MECHANISMS, BRIEFLY

Suppose that a mediator offers deals of the form $\langle x_i, n_i \rangle$, $i = w, s$, where x_i is the peace deal that will be implemented if both sides survive n_i periods of fighting. The rebel group then reports “ s ” or “ w ,” and the mediator implements $\langle x_s, n_s \rangle$ or $\langle x_w, n_w \rangle$. We also need to say what happens if the players reject the mechanism. Following the standard approach, suppose this implies “no trade,” which here means that the parties get their expected payoffs for an “absolute war” (no settlements).

We are thus assuming that the combatants (a) for some reason cannot bargain but must fight to the finish if they reject the mechanism, and (b) can commit both to fight for a given length of time and to implement any given deal, even if it is not ex post rational. Neither are remotely realistic for the anarchical context of armed conflict bargaining, but comparisons to the non-cooperative game are still useful.²⁸

For cases with $\beta_s \geq \beta_w$, it is not difficult to show that in an efficient mechanism R chooses between $\langle x_w, 0 \rangle$ and $\langle \underline{x}_s, n_s \rangle$, so that settlement occurs with no fighting when $R = R_w$ and there is an armed conflict of at most $n_s \geq 0$ periods when $R = R_s$. We have

Proposition 3. *The mechanism that maximizes the sum of ex ante expected payoffs achieves the first best of certain peace ($n_s = 0$) if either $\pi - \underline{x}_s \geq w_w$ or if the initial belief $\mu = \mu/(1 - \mu)$ is*

²⁸Assumption (a) seems quite strong even in the standard buyer-seller mechanism design setting, where contracts are feasible. Why shouldn't the parties be able to bargain with each other if they reject the mechanism?

below a threshold value defined by

$$o \leq o^M \equiv \frac{\pi - \underline{x}_s - w_s}{w_w - (\pi - \underline{x}_s)}.$$

Otherwise, if $o > o^M$, in the optimal mechanism n_s is the smallest integer such that

$$(4) \quad \frac{\mu}{1 - \mu} \leq \frac{\pi - \underline{x}_s - w_s}{\underline{x}_s - \underline{x}_w} \frac{(\delta\beta_s)^{n_s}}{(\delta\beta_w)^{n_s} - A},$$

and $x_w = (1 - (\delta\beta_w)^{n_s})\underline{x}_w + (\delta\beta_w)^{n_s}\underline{x}_s$ for this n_s .

Condition (4) comes from combining the incentive constraint for R_w and the participation constraint for G , which bind when the first-best is not possible. These are

$$(IC_w) \quad x_w \geq (1 - (\delta\beta_w)^{n_s})\underline{x}_w + (\delta\beta_w)^{n_s}\underline{x}_s,$$

$$(IR_G) \quad \mu w_w + (1 - \mu)w_s \leq \mu(\pi - x_w) + (1 - \mu)[(1 - (\delta\beta_s)^{n_s})w_s + (\delta\beta_s)^{n_s}(\pi - \underline{x}_s)]$$

$$0 \leq \mu(\pi - x_w - w_w) + (1 - \mu)(\delta\beta_s)^{n_s}(\pi - \underline{x}_s - w_s).$$

Several helpful observations follow from consideration of these conditions.

First, the central tradeoff: More war (greater n_s in the $\langle \underline{x}_s, n_s \rangle$ option) makes R_w more willing to separate itself by accepting a lower x_w . Lower x_w makes participation in the mechanism more attractive for G , but higher n_s makes it less attractive. The optimal mechanism chooses n_s (and so, through binding IC_w , x_w) to make G just willing to participate.²⁹

Second, although the mechanism has been described in dynamic terms – that is, fighting occurs for as many as n_i periods if $\langle x_i, n_i \rangle$ is chosen – it can be redescribed in static terms, akin to the approach pioneered by Myerson and Satterthwaite's (1983) analysis of bilateral trade. Replace the $(\delta\beta_i)^{n_s}$ terms in IC_w and IR_G with ϕ_w and ϕ_s , where ϕ_i is now interpreted as the probability that the mediator implements the deal $\underline{x}_s$ and does not impose permanent war (w_i for G and $\underline{x}_i$ for R) if R chooses the option for R_s .

If $\phi_w = \phi_s$, which is the case in our application when $\beta_w = \beta_s$, the weak type of rebel group can get exactly the strong type's probability of war versus settlement by mimicking the strong type. The problem is then isomorphic to the buyer-seller trade problem. By contrast, if $\beta_w < \beta_s$, then

²⁹This also implies that all the (information) rents go to R_w .

$\phi_w < \phi_s$ when $0 < n_s < \infty$. Lowering β_w makes R_w easier to separate by reducing x_w for any given $n_s \in (0, \infty)$, so relaxing G 's participation constraint, allowing smaller n_s . This captures the impact of direct learning from fighting: It is costly for the weak type to mimic the strong type not only because $\underline{x}_w$ is lower, but also because it is less likely to survive armed conflict. In other words, fighting is a kind of trial by fire that can reduce information rents.

Third, inspection shows that for the government's participation constraint to bind, it must be that $\pi - x_w - w_w < 0$. This means that the government would prefer to go to war if R chooses $\langle x_w, 0 \rangle$, revealing itself as the weak type. Thus the government *necessarily suffers ex post regret in the optimal mechanism if facing the weak type of rebel group*.

Suppose instead that we require that ex post, both parties must prefer implementing the mechanism to the “no trade” payoff of war to the finish. This is a mechanism design parallel to the “no-commitment” version of the game analyzed in the next section. The optimal mechanism with this ex post constraint has $x_w = \pi - w_w$, which gives G just enough to prefer not to fight if $R = R_w$. Then $\hat{n}_s$ is chosen as the smallest integer such that IC_w is satisfied,

$$x_w = \pi - w_w \geq (1 - (\delta\beta_w)^{\hat{n}_s})\underline{x}_w + (\delta\beta_w)^{\hat{n}_s}\underline{x}_s,$$

which we have already seen above is

$$\hat{n}_s \Delta \approx \frac{\log 1/A}{\rho + \lambda_w} = \hat{T}.$$

Inspection of condition (4) shows that with full ex ante commitment for G , the optimal maximum fighting duration n_s is strictly less than $\hat{n}_s$ ($T < \hat{T}$ for small Δ). Thus, the ex post participation constraint for G reduces aggregate welfare and increases expected fighting duration in the mechanism. We will see next that in the bargaining game the strategic pathologies implied by G 's inability to commit to deals ex post can be much greater still.

Table 2 summarizes results for different cases in both the mechanism and bargaining game analyses in the form of statements about limiting maximum war duration. Notice that when $\lambda_w = \lambda_s$, the Coase conjecture obtains in the bargaining game in exactly the same conditions as in the mechanism design setting (war duration goes to zero as the offer or battle rate gets large).

TABLE 2. Limiting maximum durations of armed conflict

Mechanism	$A \geq 1$ or $o \leq o^M$		$A < 1$ and $o > o^M$	
		$\lambda_w = \lambda_s$	$\lambda_w > \lambda_s$	
Full commitment	0	$T^M = \hat{T} + \frac{\log(1-o^*(1-A)/o)}{\rho+\lambda_w} < \hat{T}$	$0 < T^M < \hat{T}$	
With ex post participation constraint for G	0	$\hat{T}$	$\hat{T}$	

Game	$A \geq 1$ or $o \leq o^*$		$A < 1$ and $o > o^M$	
		$\lambda_w = \lambda_s$	$\lambda_w > \lambda_s$	
G can commit to deals	0	$2\hat{T}$	$0 < T^* < \hat{T}$	
G can return to war $_G$	0	∞ ; mean = $\frac{1-\mu}{1-\mu^*} \frac{1}{\lambda}$	$\frac{\log o/o^*}{\lambda_w - \lambda_s} (> T^*)$	
G can return to war $_{R_w}$	0	∞ ; mean = $\frac{1-\mu}{1-\mu^*} \frac{1}{\lambda + (\rho + \lambda)B}$	$\hat{T}$ if $o > \hat{o}$; $\frac{\log o/o^*}{\lambda_w - \lambda_s}$ if $o \in (o^*, \hat{o})$	

Notes:

- (1) $A = \frac{\pi - \underline{x}_w - w_w}{\underline{x}_s - \underline{x}_w}$ = ratio of surplus $|R_w$ to G 's max gain for separating.
- (2) $\hat{T} = -\log A/(\rho + \lambda_w)$ = fighting time s.t. R_w indiff between deal that gives G its reservation value w_w against R_w , and fighting till $\hat{T}$ to get $\underline{x}_s$.
- (3) $o^M = \frac{\pi - \underline{x}_s - w_s}{w_w - (\pi - \underline{x}_s)}$; $o^* = \frac{\rho + \lambda_s}{\rho + \lambda_w} o^M$. $\hat{o}$ is defined in Proposition 5.
- (4) In ...war $_G$, R_w gets $\underline{x}_w$ if type is revealed. In ...war $_{R_w}$, gets $\pi - w_w$.
- (5) T^* solves equation (3).
- (6) ∞ means that $Pr(\text{war duration} > T) > 0$ for all finite T . $B = \frac{\pi - \underline{x}_w - w_w}{w_w - (\pi - \underline{x}_s)}$.

When there is direct learning from fighting ($\lambda_w > \lambda_s$), there are initial beliefs $o \in (o^*, o^M)$ such that the first best of no war is possible in the mechanism but war occurs in the bargaining game. This is because the mechanism setting allows G to commit to a risk of war-to-the-finish against the strong type, which is not ex post credible and cannot occur in the non-cooperative bargaining game.

8. FIGHTING RATHER THAN BARGAINING

In the bargaining equilibrium described above, the government makes offers $x^t \in (\pi - w_w, \underline{x}_s)$ that are accepted only by the weak type of rebel group, but which are worse for the government

than permanent war with the weak type would be. The extensive form thus allows the government a power to commit that is unreasonable in a setting where no third party can reliably enforce agreements and the sides can always choose to use force.

We now consider the implications of making the more realistic assumption for this context: The government (or the offering state) cannot commit to abide by an agreement that is worse for it than fighting, given its belief about the opponent's capabilities.

Consider the following modification: In each period, if the rebel group accepts an offer x^t , the government then has another choice of whether to implement the deal or instead continue fighting. Implementation ends the game with (time averaged) payoffs $(\pi - x^t, x^t)$, while fighting yields (g, r_i) in the current period and the same continuation game as if the rebel group had rejected x^t to begin with – a lottery on military outcomes and, if stalemate, the next period. The key difference from the game with commitment is that now, if accepting an offer $x^t < \underline{x}_s$ reveals that R is the weak type, the government can go back to war if the deal is not as good as fighting with the weak type ($w_w > \pi - x^t$). We call this extensive form the *no-commitment game*.

It is straightforward to show that with complete information about rebel type $i = w, s$ the no-commitment game has a unique Markov-perfect equilibrium: The government offers $\underline{x}_i$ in every period, the rebel group always accepts any offer of at least this much, and the government always implements any x smaller than a threshold value x_i^G slightly greater than $\underline{x}_i$.

However, in contrast to the game with commitment, the no-commitment game has multiple subgame perfect Nash equilibria even with complete information. We can use the Markov perfect equilibrium, which gives the rebel group its reservation payoff, as a threat point the players expect to revert to if the rebels were to accept an offer less than some x^* , which can be greater than $\underline{x}_i$. Formally we have

Proposition 4. *In the no-commitment game with complete information and known rebel type $i = w, s$, for small enough time between offers $\Delta > 0$, there exist subgame perfect equilibria in which the government offers x^* and R_i accepts on the equilibrium path, and the government chooses to*

implement after R_i accepts. Such equilibria exist for $x^* \in \underline{x}_i \cup [x_i^G, \pi - w_i]$, where $x_i^G = \delta\beta_i\underline{x}_i + (1 - \delta\beta_i)(\pi - w_i)$.

Substantively, once the rebel's type is revealed, we can say that the two parties face a complete information bargaining problem in which the set of feasible agreements (that both sides prefer to war) is $[\underline{x}_i, \pi - w_i]$. We will analyze the no-commitment game with incomplete information for two polar cases, that differ in what the parties expect to happen in equilibrium if the weak type accepts an offer that reveals its type and the government then declines to implement. In the first case, the players expect that whenever G believes that it faces R_i for sure, they will play the Markov perfect equilibrium that implements $x = \underline{x}_i$, which is G 's favorite feasible outcome. In the second case, they expect that if R_i is revealed they will play a subgame perfect equilibrium that implements $x = \pi - w_i$, which is R_i 's favorite feasible deal.³⁰

For both cases we obtain striking results. Beginning with the first, we have

Proposition 5. *Consider updating-consistent equilibria in the no-commitment game in which $\beta_s > \beta_w > 0$ and the players expect to implement $\underline{x}_i$ in any continuation where G believes that R is certainly type R_i . For small enough time between offers $\Delta > 0$, such equilibria exist and must have the following features:*

- *On the path, after $N^* - 1$ periods of fighting (if both sides survive this long), G offers a deal x^* where $x^* \in [\underline{x}_s, \pi - w_s]$ and both types of R accept, ending the conflict. There is zero chance of a deal (an accepted, implemented offer) before period $t = N^*$. Offers before then are either rejected by both types or accepted by both types of R , but not implemented by G .*
- *Maximum war duration, $N^*\Delta$ is increasing in x^* . Thus the shortest feasible maximum conflict duration obtains for the equilibrium with $x^* = \underline{x}_s$. In this equilibrium, N^* is the*

³⁰For the case of interstate conflict, an important alternative interpretation of the second polar case ($x = \pi - w_w$ is implemented) is that both states' acquiescence is needed to change control of territory. Thus state 2 (R) can agree to cede more than w_w , but in the complete information situation state 1 cannot credibly threaten to go to war rather than live with any territorial division $x \leq \pi - w_w$.

smallest integer such that

$$N^* \Delta \geq \frac{\log o/o^*}{\lambda_w - \lambda_s},$$

where o^ is the same belief threshold defined in Proposition 2. This is strictly greater than the limiting maximum duration in the game with commitment, and grows without bound in the initial belief o that R is the weak type of rebel group.*

In words, we immediately get pure *fighting rather than bargaining*. The government makes non-serious offers – offers that all know have zero probability of acceptance or implementation – until its belief that R is the strong type increases to the point that it is willing to make the pooling offer, $x^* \geq \underline{x}_s$. Acceptance of this offer (by both types) leads to a self-enforcing peace.

The logic is straightforward. If the weak type reveals itself by accepting an offer the strong type would certainly reject, its continuation payoff if G goes back to war is exactly its reservation value for “pure war,” $\underline{x}_w$, because it has no bargaining power in the continuation game and will be pushed by G to this limit. But by rejecting offers less than x^* , *and so maintaining the reputation of possibly being the strong type*, it gets its war payoff while fighting ($\underline{x}_w$) followed by a strictly better deal, x^* , if it can manage to survive long enough on the battlefield, which has positive probability as long as $\beta_w > 0$.

“Fighting rather than bargaining” can dramatically increase expected and maximum conflict durations. Maximum duration in this no-commitment case is the time it takes for pure battlefield learning (the mechanical effect alone) to decrease G ’s optimism to the point that it is willing to make the offer a strong type would accept. The minimum such duration, $T_{nc} = \log(o/o^*)/(\lambda_w - \lambda_s)$, is of course strictly greater than, T^* , the limiting maximum on war duration in the commitment case with $\lambda_w > \lambda_s$. T_{nc} can be in years for reasonable parameter values. Note also that T_{nc} increases without bound in the initial odds that $R = R_w$. (By contrast, $\hat{T}$ is the upper bound for any initial belief in the game with commitment to accepted offers.)

In strategic terms, fighting-rather-than-bargaining and possibly very long war durations are driven by a “ratchet effect” like that noted in models of bargaining over a rental good or service.

Such goods entail flow payoffs (as in the armed conflict setting), which raise the concern that accepting today might undermine leverage in renegotiation tomorrow by revealing that one is a higher-valuation type than was previously believed. In the buyer-seller context with independent values, a seller's inability to commit not to exploit a buyer who says Yes today ends up producing *immediate* trade – no delay, so the Coase conjecture obtains despite lack of commitment to offers. In that setting, the ratchet effect confers all bargaining power on the buyer (informed party), because the seller's inability to commit to her proposal means that the buyer has no incentive to accept any screening offer. As a result the seller's best option is simply to make the pooling offer right away.³¹

By contrast, in the armed conflict setting when $\beta_s > \beta_w$, the government has an alternative to screening via offers – *screening by fighting*. Due to interdependent values about the “no trade” (fighting) option, the government can learn by putting the rebel group through trial by fire. The government's belief that the rebels are the tough type rises the longer they survive purely due to the mechanical effect, eventually to the point that the government is willing to make the better pooling offer to end the war on terms acceptable to both rebel types.

This stark outcome – pure fighting rather than bargaining when $o > o^*$ – is clearly related to the government having all the bargaining power if a weak type reveals itself. In effect this limits the set of feasible agreements to $\underline{x}_w$ and $x^* \in [\underline{x}_s, \pi - w_s)$ since the strong type will never accept a deal worse than $\underline{x}_s$ and the weak type can't expect anything better than $\underline{x}_w$ due to the government's inability to commit not to exploit its bargaining power in complete-information renegotiations. To the combatants and outside observers it may appear as if the issues in question are “indivisible” and in effect they are. But this is a strategic consequence rather than an inherent property of the issues (cf. Powell 2006).

We next show that fighting rather than bargaining remains even if the government does not have all bargaining power when rebel type is known.

³¹See Hart and Tirole (1988), the case of short-term contracts. Their setting is independent values, so that the seller gets a surplus from trading with a high valuation buyer even at the low valuation buyer's reservation price. As far as we know our model and analysis are the first of bargaining over a flow payoff with interdependent values and no commitment to implement accepted deals.

Proposition 6. *Consider the no-commitment game where $A < 1$ and $\beta_s > \beta_w > 0$. Restrict attention to equilibria that satisfy updating consistency and in which the players expect to implement $\pi - \underline{x}_w$ in any continuation where G believes that R is certainly the weak type, and G never offers more than $\underline{x}_s$. For small enough time between offers $\Delta > 0$ there exists a $\hat{\mu} > \mu^*$ such that*

- (1) *for initial beliefs $\mu \in (\mu^*, \hat{\mu})$, any such equilibrium is the same as described in Proposition 5 (with $x^* = \underline{x}_s$);*
- (2) *for $\mu \in (\hat{\mu}, 1)$, any such equilibrium follows the same pattern as in Proposition 1 (the commitment case) for periods $t = 0, 1, \dots, N - \hat{n}$, and then switches to the pattern of Proposition 5 (N is maximum duration less one). $\hat{n}$ is the smallest integer n such that $\pi - w_w \geq ((1 - (\delta\beta_w)^n)\underline{x}_w + (\delta\beta_w)^n\underline{x}_s$; in the limit $\hat{n}\Delta = \hat{T} = \log(1/A)/(\rho + \lambda_w)$. Thus, G initially makes ascending offers $x^t = x_{N-t}$ which R_s rejects for sure while R_w rejects with a positive probability ν^t until reaching the ex post regret point of $x_{\hat{n}}$. After this G makes any non-serious offers less than $\underline{x}_s$ that R rejects for sure until fighting causes G 's belief to fall to $\mu^t \leq \mu^*$, when G makes the pooling offer $\underline{x}_s$.*

Further $\lim_{\Delta \downarrow 0} \hat{\mu} = \mu^*/(\mu^* + A^{L/\rho'}(1 - \mu^*))$.

In this case, if R_w accepts an offer $x^t < \underline{x}_s$ that reveals its type, G is then choosing between implementing x^t and continuing the war. The latter would lead in the next period to a settlement at G 's reservation for war with the weak type, $\pi - w_w$.³² Here, if the initial belief that R is the weak type is large enough, then the conflict begins with an increasing sequence of offers $x^t < \pi - w_w$ that may be accepted by R_w , and which would then be implemented by G . But there is positive probability that even the weak type will reject all these offers and the conflict will enter a phase of fighting rather than bargaining. Once x^t rises to $\pi - w_w$, the government cannot commit to

³²For the interstate war context, this case has the following natural interpretation: State 1 (G) is simply making verbal proposals, and both states must agree to implement a change in the territorial status quo. Once state 2 (R_w) reveals its type, it can cede territory just up to $\pi - w_w$, at which point state 1 no longer has a credible threat to fight. Whereas Fearon (1995) and similar models assume that state 1's "demand" is a unilateral change in the status quo that then confronts the other with a choice of acquiescence or war, Powell (1996b,a, 2004) analyzes models in which both states must agree to change a status quo division. In these, a status quo is stable if it is common knowledge that both prefer it to war.

implement more than this (but less than $\underline{x}_s$), and R_w strictly prefers, in equilibrium, to mimic the strong type in hopes of surviving long enough to get $\underline{x}_s$.

In the limit as the time between offers gets small, the initial bargaining phase happens very quickly, tending to instantaneous. Maximum war duration is thus larger than in the commitment case – it tends to $\hat{T}$ for any $\mu > \hat{\mu}$, and is strictly greater than $\hat{T}$ for $\mu \in (\mu^*, \hat{\mu})$. Again, G 's inability to commit to implement deals that imply ex post regret increases inefficiency by undermining screening through offers.

9. CONCLUSION

Several basic empirical features of interstate and civil war are hard to reconcile with standard models of buyer-seller bargaining, despite good reasons to conceive of war as some kind of bargaining process. First, war durations, particularly for civil wars, can be remarkably long despite high levels destruction and costs often born while fighting. In the post-World War II period, median and maximum duration for interstate wars were 3.4 months and about 10 years, respectively. For civil wars the corresponding median and maximum duration were eight years and more than 60 years.³³

Second, combatants usually declare and stick with extreme war aims that no one expects the other side to concede except after total military defeat. They choose to *fight rather than to bargain*, if by bargaining we mean an interaction involving the constant reformulation and exchange of serious offers. Third, both interstate and civil wars sometimes do end with negotiated settlements that are the product of serious offers at a bargaining table, but such bargaining happens quickly and tends to come after significant periods of fighting without offers that have positive probability of acceptance in the near term. Fourth, historians and political scientists have argued for many cases that the reason that serious bargaining can finally begin is that one or both sides have revised downwards their beliefs about the prospects of military victory based on the accumulation of battlefield evidence.³⁴

³³Using the Correlates of War data for interstate wars, and Fearon (2004) for civil wars.

³⁴See especially Blainey (1973) and the large literature following him; also Zartman (1989).

This paper clarifies the relationship between buyer-seller and armed conflict bargaining. In the latter, the “good” can change hands not only via a deal but also by military victory. We find that with private information about military prospects, armed conflict bargaining is analogous to buyer-seller bargaining with interdependent values, but differs in that the private information concerns the value of “no trade” – here, fighting – rather than the value of good at issue (territory, foreign policies, etc.). That fighting can provide direct information about a combatant’s future prospects then affects negotiating positions, in a way we find makes for a more slow but steady change in beliefs about the other side’s capabilities in the early phases of a war.

We also find that when we introduce the natural assumption that the parties can’t commit not to change a deal after an agreement, we immediately get fighting rather bargaining in equilibrium, along with negotiated settlements that ultimately become possible only because of mechanical (i.e., battlefield) learning about the balance of military power. In anarchical settings, states or rebel groups can reasonably fear that accepting a deal will lead the other side to conclude that they can be pushed even further. This problem undermines the parties’ ability to discriminate between tougher and weaker adversaries through the negotiating process, leaving fighting without making serious offers as a next best alternative for screening between types. It also can greatly increase the median and maximum durations of a conflict.

There are two main alternative explanations for “fighting rather than bargaining.”³⁵ First, in some models protracted fighting without serious offers is driven by a commitment problem occasioned by the fear that military power would shift against one side if it fails to fight or stops fighting (e.g., Fearon, 1998, 2004; Powell, 2006). Second, since Blainey (1973) a common argument has been that wars occur when the complexity of making military estimates allows two boundedly

³⁵Far less developed are models that would explain fighting rather than bargaining by reference to principal-agent problems in domestic politics. Goemans (2000) argues that the German government in World War I held their war aims steady or even increased them after some battlefield defeats because they expected to be deposed in a revolution if their war gains were not very large; they were “gambling for resurrection.” Some models of gambling for resurrection have been developed for decisions to go to war (e.g., Downs and Rocke, 1994; Hess and Orphanides, 1995) but work extending the idea to intra-war bargaining is in its infancy.

rational or otherwise biased leaderships to be mutually optimistic about their odds of winning (e.g., Wagner, 2000; Smith and Stam, 2004). War gradually teaches them who is right.

Both can be plausible for particular cases, and there is no reason why all three mechanisms might not operate in the same case.³⁶ Nonetheless, two empirical observations suggest that the ratchet effect problem modeled here may have added purchase. First, the shifting-power commitment problem explanation implies wars to the finish – no negotiated settlements. As noted, however, both civil and interstate wars sometimes end short of a decisive military resolution, with a negotiated settlement.³⁷

Second, the bounded-rationality/mutual optimism explanation predicts that combatants will *continuously* adjust their war aims and settlement offers as they update about relative military capabilities (Wagner, 2000; Goemans, 2000; Reiter, 2009). But states' and rebel groups' war aims in conflicts tend to be very sticky, experiencing little change until there is a sudden shift into negotiations (Iklé, 1971). A main finding of Reiter's (2009) examination of 22 "key decision" times concerning termination in six major interstate wars is that often state leaders did not change their aims or demands *at all*, or even increased them, after major battlefield setbacks. He gives "fears of adversarial non-compliance with war-ending agreements" and "the fear decision-makers sometimes have that making concessions will send signals of weakness" as two of three main reasons.³⁸

From a normative perspective it could be useful to know whether protracted fighting due the ratchet effect or a pure commitment problem is more or less empirically common than protracted fighting due to ill-founded mutual optimism, and also how to identify one or the other in particular

³⁶For example, Wolford, Reiter and Carrubba (2011) analyze a model in which one state is uncertain about whether an adversary's costs for fighting are low enough that there will be no feasible deal in the future due to a shifting-power commitment problem.

³⁷The very long durations of many civil wars, and other features of civil wars, suggests strongly that they are often driven by the problem of commitment to powersharing deals. For example, they are much less likely to end with negotiated settlements than are interstate wars. See Fearon (1998, 2004); Walter (1997).

³⁸Reiter 2009, 221. His third main reason is that "leaders can be patient" and hope to turn things around in the future; this can result from "overconfidence" but also might reflect, as in the model, the weak type's incentive to hang on for the good offer even if the odds of getting it are low. Note that in Reiter's cases "fears of adversarial non-compliance" arise both due to fears of power shifts and fears of loss of reputation (as in "signals of weakness"). Pillar (1983), Iklé (1971), and Smith (1995, chapter 3) also stress fear of appearing weak as a significant barrier to serious negotiations in interstate wars.

cases. When the problem is primarily mutual optimism, the most relevant policy intervention may be third-party provision of better military analysis in hopes of bringing expectations into accord. If the problem is the ratchet effect or the fear that bargaining power will shift for some other reason in the future, then the most relevant policy interventions may be third-party efforts to guarantee and enforce a negotiated deal.³⁹

In this paper we did not consider the case of $\beta_w > \beta_s$, which would describe a conflict where the rebel group's failure to win decisively as time passed would tend to indicate that it was more likely the weak type. We conjecture that in this case pure fighting rather than bargaining could result in equilibrium even if the government has the ability to commit to offers. It would also be of interest to explore armed conflict bargaining with private information about military capabilities when it is assumed, or allowed, that the informed party makes offers. It is immediately clear that in contrast to the independent values case, inefficient delay could still be necessary in equilibrium, despite the informed party having all the structural bargaining power (as shown in Gerardi, Hörner and Maestri (2014), who study this case for the buyer-seller problem).

³⁹See, for example, Walter (2002) on third-party guarantees as an important component of civil war settlements.

REFERENCES

- Admati, Anat R. and Motty Perry. 1987. "Strategic Delay in Bargaining." *Review of Economic Studies* 54(3):345–364.
- Akerlof, George A. 1970. "The Market for "Lemons": Quality Uncertainty and the Market Mechanism." *Quarterly Journal of Economics* 84(3):488–500.
- Ausubel, Lawrence M., Peter Cramton and Raymond J. Deneckere. 2002. Bargaining with Incomplete Information. In *Handbook of Game Theory, Vol. 3*, ed. Robert J. Aumann and Sergiu Hart. Amsterdam: Elsevier Science B.V. chapter 50.
- Bester, Helmut and Karl Wärneryd. 2006. "Conflict and the Social Contract." *Scandinavian Journal of Economics* 108(2):231–249.
- Blainey, Geoffrey. 1973. *The Causes of War*. New York: Free Press.
- Card, David. 1990. "Strikes and Wages: A Test of an Asymmetric Information Model." *Quarterly Journal of Economics* 105(3).
- Coase, Ronald H. 1972. "Durability and Monopoly." *Journal of Law and Economics* 15:143–149.
- Deneckere, Raymond and Meng-Yu Liang. 2006. "Bargaining with Interdependent Values." *Econometrica* 74(5):1309–1364.
- Downs, George W. and David M. Rocke. 1994. "Conflict, Agency, and Gambling for Resurrection." *American Journal of Political Science* 38(2):362–80.
- Fearon, James D. 1995. "Rationalist Explanations for War." *International Organization* 49(3):379–414.
- Fearon, James D. 1998. Commitment Problems and the Spread of Ethnic Conflict. In *The International Spread of Ethnic Conflict*, ed. David A. Lake and Donald Rothchild. Princeton, NJ: Princeton University Press pp. 107–26.
- Fearon, James D. 2004. "Why Do Some Civil Wars Last So Much Longer Than Others?" *Journal of Peace Research* 41(3):275–301.
- Fearon, James D. 2013. "Fighting rather than Bargaining." Unpublished paper.
- Fey, Mark, Adam Meirowitz and Kristopher W Ramsay. 2013. "Credibility and Commitment in Crisis Bargaining." *Political Science Research and Methods* 1(1):27–52.
- Fey, Mark and Kristopher W. Ramsay. 2011. "Uncertainty and Incentives in Crisis Bargaining: Game-Free Analysis of International Conflict." *American Journal of Political Science* 55(1):149–169.
- Filson, Darren and Suzanne Werner. 2002. "A bargaining model of war and peace: Anticipating the onset, duration, and outcome of war." *American Journal of Political Science* pp. 819–837.
- Forges, Francoise. 1999. Ex post individually rational trading mechanisms. In *Current Trends in Economics*, ed. Ahmet Alkan, Charalambos D. Aliprantis and Nicholas C. Yanneli. Springer pp. 157–175.
- Fudenberg, Drew and Jean Tirole. 1991. *Game Theory*. Cambridge, MA: MIT Press.
- Gerardi, Dino, Johannes Hörner and Lucas Maestri. 2014. "The role of commitment in bilateral trade." *Journal of Economic Theory* 154:578–603.
- Goemans, Hein. 2000. *War and Punishment*. Princeton, NJ: Princeton University Press.
- Gul, Faruk, Hugo Sonnenschein and Robert Wilson. 1986. "Foundations of Dynamic Monopoly and the Coase Conjecture." *Journal of Economic Theory* 39(1):155–190.
- Hart, Oliver D. 1989. "Bargaining and Strikes." *Quarterly Journal of Economics* 104(1):25–43.

- Hart, Oliver D. and Jean Tirole. 1988. "Contract Renegotiation and Coasian Dynamics." *Review of Economic Studies* 40:509–40.
- Hess, Gregory D. and Athanasios Orphanides. 1995. "War Politics: An Economic, Rational-Voter Framework." *American Economic Review* 85(4):828–46.
- Iklé, Fred Charles. 1971. *Every War Must End*. New York: Columbia University Press.
- Min, Eric. 2017. "Negotiation as an Instrument of War." Stanford University.
- Myerson, Roger B. and Mark Satterthwaite. 1983. "Efficient Mechanisms for Bilateral Trading." *Journal of Economic Theory* 29:265–281.
- Perea, Andrés. 2002. "A Note on the One-Deviation Property in Extensive Form Games." *Games and Economic Behavior* 40:332–338.
- Pillar, Paul R. 1983. *Negotiating Peace: War Termination as a Bargaining Process*. Princeton, NJ: Princeton University Press.
- Powell, Robert. 1996a. "Bargaining in the Shadow of Power." *Games and Economic Behavior* 15:255–89.
- Powell, Robert. 1996b. "Stability and the Distribution of Power." *World Politics* 48(2):239–267.
- Powell, Robert. 2002. "Bargaining Theory and International Conflict." *Annual Review of Political Science* 5:1–30.
- Powell, Robert. 2004. "Bargaining and Learning While Fighting." *American Journal of Political Science* 48(2):344–361.
- Powell, Robert. 2006. "War as a Commitment Problem." *International Organization* 60(1):169–204.
- Reiter, Dan. 2003. "Exploring the Bargaining Model of War." *Perspectives on Politics* 1(1):27–43.
- Reiter, Dan. 2009. *How Wars End*. Princeton: Princeton University Press.
- Shimer, Robert and Iván Werning. 2015. Efficiency and information transmission in bilateral trading. Technical report National Bureau of Economic Research.
- Slantchev, Branislav L. 2003. "The Principle of Convergence in Wartime Negotiations." *American Political Science Review* 97(1):621–32.
- Smith, Alastair and Allan C. Stam. 2004. "Bargaining and the Nature of War." *Journal of Conflict Resolution* 48(6):783–813.
- Smith, James D. 1995. *Stopping Wars: Defining Obstacles to a Cease-Fire*. Boulder, CO: Westview Press.
- Wagner, R. Harrison. 2000. "Bargaining and War." *American Journal of Political Science* 44(3):469–84.
- Walter, Barbara. 2002. *Committing to Peace*. Princeton: Princeton University Press.
- Walter, Barbara F. 1997. "The Critical Barrier to Civil War Settlement." *International Organization* 51(3):335–64.
- Wittman, Donald. 2009. "Bargaining in the shadow of war: when is a peaceful resolution most likely?" *American journal of political science* 53(3):588–602.
- Wolford, Scott, Dan Reiter and Clifford J. Carrubba. 2011. "Information, commitment, and war." *Journal of Conflict Resolution* 55(4):556–579.
- Zartman, I. William. 1989. *Ripe for Resolution: Conflict and Intervention in Africa*. Oxford: Oxford University Press.

Appendix

1. PROPOSITION 1

1.1. Preliminaries for proof of Proposition 1. We begin by introducing notation and deriving the recursion that produces the equilibrium path of G 's beliefs and R_w 's mixing probabilities. Then the proof below first demonstrates uniqueness (generically with respect to the initial belief), and next checks that it forms an updating-consistent PBE.

First, some notation.

- $\psi(\mu) = \frac{\mu\beta_s}{(1-\mu)\beta_w + \mu\beta_s}$. $\psi(\mu)$ is the prior on R_w such that after one round of pure fight the posterior is exactly μ . (Thus $\psi^{-1}(\mu)$ is Bayes' rule with prior μ .) We assume throughout that $\beta_s \geq \beta_w$, which implies that $\psi(\mu) \geq \mu$.
- $o(x) = \frac{x}{1-x}$.
- $\tilde{g}_i = g(1-\delta) + \delta\gamma_i\pi$. Note that $\tilde{g}_i = (1-\delta\beta_i)w_i$.
- $\bar{v} = \pi - g - r_w$.
- $\underline{v} = \pi - g - r_s$.

Consider a sequence of beliefs $\{\mu_n\}_{n=0}^\infty$ with $\mu_0 = 0$ and $\mu_n \geq \psi(\mu_{n-1})$ for $n \geq 1$. Let

$$\nu_n(\mu) \triangleq \frac{o(\mu_{n-1})\beta_s}{o(\mu)\beta_w}.$$

Substantively, $\nu_n(\mu)$ is R_w 's probability of rejecting an offer such that G would update from μ to μ_{n-1} on seeing rejection (and assuming R_s rejects for sure).

For $n \geq 1$ and $\mu \geq \psi(\mu_{n-1})$, recursively define the functions

$$\begin{aligned} u_0^w(\mu) &\equiv \pi - \underline{x}_s \\ u_n^w(\mu) &\triangleq (1 - \nu_n(\mu))(\pi - x_n) + \nu_n(\mu)(\tilde{g}_w + \delta\beta_w u_{n-1}^w(\mu_{n-1})) \\ u_n^s &\triangleq \tilde{g}_s \frac{1 - (\delta\beta_s)^n}{1 - \delta\beta_s} + (\delta\beta_s)^n(\pi - \underline{x}_s) = (1 - (\delta\beta_s)^n)w_s + (\delta\beta_s)^n(\pi - \underline{x}_s) \\ u_n(\mu) &\triangleq \mu u_n^w(\mu) + (1 - \mu)u_n^s. \end{aligned}$$

$u_n^i(\mu)$ is G 's expected payoff conditional on facing R_i with $n+1$ periods till bargaining must end, given belief μ . $u_n(\mu)$ is then the expected total payoff for G , given initial belief μ , if it chooses the bargaining process that lasts at most $n+1$ rounds given the sequence $\mu_{n-1}, \mu_{n-2}, \dots, \mu_0$. $u_{n-1}(\mu)$ is the payoff if it chooses the one that lasts at most n rounds given μ .

Next, derivation of the recursion.

First we show that the unique sequence $\{\mu_n\}_{n=0}^{\infty}$ that satisfies $\mu_0 = 0$ and $u_n(\mu_n) = u_{n-1}(\mu_n)$ for $n \geq 1$ is given by the recursion (2). In words, $u_n(\mu_n)$ is G 's equilibrium path payoff with n periods to go and $u_{n-1}(\mu_n)$ is G 's payoffs to "skipping ahead" by making the offer x_{n-1} given belief μ_n . Writing these out we have

$$u_n(\mu_n) = \mu_n [(1 - \nu_n)(\pi - x_n) + \nu_n(\tilde{g}_w + \delta\beta_w u_{n-1}^w)] + (1 - \mu_n) \left[\tilde{g}_s \frac{1 - (\delta\beta_s)^n}{1 - \delta\beta_s} + (\delta\beta_s)^n(\pi - \underline{x}_s) \right].$$

and

$$u_{n-1}(\mu_n) = \mu_n [(1 - \nu'_n)(\pi - x_{n-1}) + \nu'_n(\tilde{g}_w + \delta\beta_w u_{n-2}^w)] + (1 - \mu_n) \left[\tilde{g}_s \frac{1 - (\delta\beta_s)^{n-1}}{1 - \delta\beta_s} + (\delta\beta_s)^{n-1}(\pi - \underline{x}_s) \right]$$

where $\nu_n = o_{n-1}\beta_s/o_n\beta_w$ and $\nu'_n = o_{n-2}\beta_s/o_n\beta_w$.

Setting $u_n(\mu_n) = u_{n-1}(\mu_n)$ and extensive algebra then leads to

$$\begin{aligned} (\delta\beta_s)^{n-1}(1 - \delta)\bar{v} &= o_n(x_{n-1} - x_n) + \frac{\beta_s}{\beta_w} o_{n-2}(\pi - x_{n-1} - \tilde{g}_w - \delta\beta_w u_{n-2}^w) \\ (5) \quad &\quad - \frac{\beta_s}{\beta_w} o_{n-1}(\pi - x_n - \tilde{g}_w - \delta\beta_w u_{n-1}^w) \\ &= o_n(x_{n-1} - x_n) + \frac{\beta_s}{\beta_w} o_{n-2}z_{n-1} - \frac{\beta_s}{\beta_w} o_{n-1}z_n \end{aligned}$$

where $z_n \equiv \pi - x_n - \tilde{g}_w - \delta\beta_w u_{n-1}^w$. (z_n is G 's gain if facing the weak type if the weak type accepts in n versus rejecting in n .)

On the other hand, note that $u_{n-1}^w = \pi - x_{n-1} - \nu_{n-1}(\pi - x_{n-1} - \tilde{g}_w - \delta\beta_w u_{n-2}^w) = \pi - x_{n-1} - \nu_{n-1}z_{n-1}$, from which we can get

$$\begin{aligned} z_n &= \pi - x_n - \tilde{g}_w - \delta\beta_w(\pi - x_{n-1} - \nu_{n-1}z_{n-1}) \\ &= (1 - \delta)(\pi - r_w - g) + \delta\beta_w\nu_{n-1}z_{n-1} \\ &= (1 - \delta)\bar{v} + \delta\beta_w\nu_{n-1}z_{n-1}, \end{aligned}$$

or

$$o_{n-1}z_n - \delta\beta_s o_{n-2}z_{n-1} = o_{n-1}(1 - \delta)\bar{v}.$$

From (5), we have

$$o_n(x_{n-1} - x_n) + \frac{\beta_s}{\beta_w} o_{n-2}z_{n-1} - \frac{\beta_s}{\beta_w} o_{n-1}z_n = \delta\beta_s \left(o_{n-1}(x_{n-2} - x_{n-1}) + \frac{\beta_s}{\beta_w} o_{n-3}z_{n-2} - \frac{\beta_s}{\beta_w} o_{n-2}z_{n-1} \right)$$

Combining the last two equations gives a second-order linear difference equation (for $n \geq 2$)

$$\begin{aligned} o_n &= \frac{1}{x_{n-1} - x_n} \left(\left(\delta\beta_s(x_{n-2} - x_{n-1}) + \frac{\beta_s}{\beta_w}(1 - \delta)\bar{v} \right) o_{n-1} - \frac{\beta_s}{\beta_w}(1 - \delta)\bar{v}o_{n-2} \right) \\ &= o_{n-1} \left(\frac{\beta_s}{\beta_w} + \frac{\beta_s}{\beta_w}A_n \right) - o_{n-2} \frac{\beta_s}{\beta_w}A_n, \end{aligned}$$

where

$$A_n = \frac{1}{(\delta\beta_w)^{n-1}} \frac{(1 - \delta)(\pi - r_w - g)}{(1 - \delta\beta_w)(\underline{x}_s - \underline{x}_w)}.$$

1.2. Proof of Proposition 1. Throughout the proof we implicitly impose the restriction of updating consistency. Recall that updating consistency requires upon deviation the beliefs be recomputed under the new strategies. Hence in what follows, whenever we consider a candidate strategy, we shall assume the beliefs are consistent with the strategy regardless of whether we are on the equilibrium path or not.

We prove the proposition in three steps.

Step 1: Uniqueness of the Equilibrium Path.

Any PBE satisfying updating consistency must have the equilibrium path specified in the proposition.

Lemma 1. *For any $n \geq 1$,*

$$\pi - x_n > \tilde{g}_w + \delta\beta_w(\pi - x_{n-1}).$$

Proof. Using the definition of x_n ,

$$\pi - x_n - (\tilde{g}_w + \delta\beta_w(\pi - x_{n-1})) = (1 - \delta)(\pi - r_w - g) > 0.$$

□

Lemma 1 says simply that on the offer path that makes R_w indifferent between accepting in n versus $n - 1$, G does better with acceptance in n versus fighting and then acceptance in $n - 1$.

Lemma 2. *Suppose $\mu \geq \psi(\mu_n)$, and thus $\psi^{-1}(\mu) \geq \mu_n$, for some $n \geq 0$. Then*

$$(6) \quad (1 - \nu)(\pi - x_{n+1}) + \nu(\tilde{g}_w + \delta\beta_w u_n^w(\mu')),$$

where $\nu = o(\mu')\beta_s/o(\mu)\beta_w$, is strictly decreasing in $\mu' \in [\mu_n, \psi^{-1}(\mu)]$.

Proof. For this proof, we always assume $\mu' \in [\mu_n, \psi^{-1}(\mu)]$. Also, let $h(\mu')$ denote the expression (6).

If $n = 0$,

$$h(\mu') = (1 - \nu)(\pi - x_1) + \nu(\tilde{g}_w + \delta\beta_w(\pi - \underline{x}_s)).$$

Since ν is strictly increasing in μ' and $h(\mu')$ is strictly decreasing in ν , $h(\mu')$ is strictly decreasing in μ' .

Now suppose $n \geq 1$. Define $\nu' \triangleq o(\mu_{n-1})\beta_s / (o(\mu')\beta_w)$. We can rewrite $h(\mu')$ as

$$\begin{aligned} h(\mu') &= (1 - \nu)(\pi - x_{n+1}) + \nu (\tilde{g}_w + \delta\beta_w ((1 - \nu')(\pi - x_n) + \nu' (\tilde{g}_w + \delta\beta_w u_{n-1}^w(\mu_{n-1}))) \\ &= \pi - x_{n+1} + \nu (\tilde{g}_w + \delta\beta_w(\pi - x_n) - (\pi - x_{n+1})) + \nu\nu'\delta\beta_w (\tilde{g}_w + \delta\beta_w u_{n-1}^w(\mu_{n-1}) - (\pi - x_n)). \end{aligned}$$

Note that $\nu\nu' = o(\mu_{n-1})\beta_s^2 / (o(\mu)\beta_w^2)$ is independent of μ' . So by Lemma (1), $h(\mu')$ is strictly decreasing in μ' . $\square$

Lemma 2 says that the higher G 's post-rejection belief that R_w , the lower G 's expected payoff per rejection.

Corollary 1. *Suppose $\mu > \psi(\mu_n)$ for some $n \geq 0$. Then*

$$u_{n+1}^w(\mu) > \tilde{g}_w + \delta\beta_w u_n^w(\psi^{-1}(\mu)).$$

And thus

$$u_{n+1}(\mu) > \mu (\tilde{g}_w + \delta\beta_w u_n^w(\psi^{-1}(\mu))) + (1 - \mu)u_{n+1}^s.$$

Proof. $u_{n+1}^w(\mu)$ is what we would get by plugging $\mu' = \mu_n$ into (6), and $\tilde{g}_w + \delta\beta_w u_n^w(\psi^{-1}(\mu))$ is what we would get by plugging $\mu' = \psi^{-1}(\mu)$ into (6). So by Lemma (2), we have

$$u_{n+1}^w(\mu) > \tilde{g}_w + \delta\beta_w u_n^w(\psi^{-1}(\mu)).$$

$\square$

In words, Corollary (1) means that if $\mu > \psi(\mu_n)$, making a serious offer that has positive chance of acceptance (following the bargaining process corresponding to u_{n+1}) is better than wasting such a chance by making an offer that induces certain rejection (which would mean that any belief change from the current period is purely driven by fighting).

Lemma 3. *For any $n \geq 1$, $u_n(\mu) > u_{n-1}(\mu)$ for $\mu \in (\mu_n, 1)$, and $u_n(\mu) < u_{n-1}(\mu)$ for $\mu \in [\psi(\mu_{n-1}), \mu_n)$.*

Proof. From the definition of u_n^w we have

$$\begin{aligned} u_n^w(\mu) &\triangleq (1 - \nu_n(\mu))(\pi - x_n) + \nu_n(\mu)(\tilde{g}_w + \delta\beta_w u_{n-1}^w(\mu_{n-1})) \\ &= \pi - x_n - \frac{o(\mu_{n-1})}{o(\mu)} \frac{\beta_s}{\beta_w} (\pi - x_n - \tilde{g}_w - \delta\beta_w u_{n-1}^w(\mu_{n-1})) \\ &= \pi - x_n - \frac{1 - \mu}{\mu} h_n, \end{aligned}$$

where we have defined, for each $n \geq 1$,

$$h_n \triangleq o(\mu_{n-1}) \frac{\beta_s}{\beta_w} (\pi - x_n - \tilde{g}_w - \delta\beta_w u_{n-1}^w(\mu_{n-1})).$$

Thus

$$\begin{aligned}\mu u_n^w(\mu) &= \mu(\pi - x_n) + (\mu - 1)h_n \\ u_n(\mu) &= \mu u_n^w(\mu) + (1 - \mu)u_n^s = \mu(\pi - x_n) + (\mu - 1)h_n + (1 - \mu)u_n^s \\ u_n(\mu) &= (\pi - x_n + h_n - u_n^s)\mu - h_n + u_n^s.\end{aligned}$$

This shows that u_n is a linear function in μ . Since we know $u_n(\mu_n) = u_{n-1}(\mu_n)$, it then suffices to show the slope of u_n is larger than the slope of u_{n-1} .

Starting with the equation $u_n(\mu_n) = u_{n-1}(\mu_n)$, we can get

$$h_n - h_{n-1} = u_n^s - u_{n-1}^s + \frac{\mu_n}{1 - \mu_n}(x_{n-1} - x_n).$$

Thus the slope difference between u_n and u_{n-1} is

$$\begin{aligned}(\pi - x_n + h_n - u_n^s) - (\pi - x_{n-1} + h_{n-1} - u_{n-1}^s) &= h_n - h_{n-1} + (x_{n-1} - x_n) + (u_{n-1}^s - u_n^s) \\ &= \frac{1}{1 - \mu_n}(x_{n-1} - x_n) \\ &> 0,\end{aligned}$$

as desired. □

Lemma 4. *If $\mu \in [\mu_n, \mu_{n+1})$ for some n , then*

$$(7) \quad u_n(\mu) > \mu w_w + (1 - \mu)w_s.$$

Proof. Let the function $h : [0, 1] \rightarrow \mathbb{R}$ be given by

$$h(\mu) = \mu w_w + (1 - \mu)w_s.$$

Note that

$$h(\mu) = \mu \tilde{g}_w + (1 - \mu)\tilde{g}_s + \delta(\mu\beta_w + (1 - \mu)\beta_s)h(\phi^{-1}(\mu)),$$

for all $\mu \in [0, 1)$.

First assume $\beta_s > \beta_w$ so that $\mu > \phi(\mu)$ for all $\mu > 0$. We proceed by induction on n for this case. The base case where $n = 0$ can be verified by direct computation.

Now suppose (7) holds for $n = k$, and we would like to show it for $n = k + 1$. Pick an $\epsilon > 0$ so that $\mu - \phi^{-1}(\mu) > \epsilon$ for all $\mu \in [\mu_{k+1}, \mu_{k+2})$. And for convenience assume $N_1 \triangleq (\mu_{k+2} - \mu_{k+1})/\epsilon$ is an integer. If $\mu_{k+1} \leq \mu < \mu_{k+1} + \epsilon$, then we must have $\mu_k \leq \phi^{-1}(\mu) < \mu_{k+1}$, and thus Corollary 1 implies

$$\begin{aligned}(8) \quad u_{k+1}(\mu) &> \mu(\tilde{g}_w + \delta\beta_w u_k^w(\phi^{-1}(\mu))) + (1 - \mu)u_{k+1}^s \\ &= \mu\tilde{g}_w + (1 - \mu)\tilde{g}_s + \delta(\mu\beta_w + (1 - \mu)\beta_s)u_k(\phi^{-1}(\mu)) \\ &> \mu\tilde{g}_w + (1 - \mu)\tilde{g}_s + \delta(\mu\beta_w + (1 - \mu)\beta_s)h(\phi^{-1}(\mu)) \\ &= h(\mu).\end{aligned}$$

Assume we have proved the situation for all $\mu \in [\mu_{k+1}, \mu_{k+1} + j\epsilon)$, $1 \leq j < N_1$. Consider the case $\mu \in [\mu_{k+1} + j\epsilon, \mu_{k+1} + (j+1)\epsilon)$. If $\mu < \phi(\mu_{k+1})$, the same logic of (8) still applies. And if $\mu \geq \phi(\mu_{k+1})$, recall from the construction of the cutoffs $\{\mu_j\}_{j=0}^\infty$ that $u_{k+1}(\mu) > u_{k+2}(\mu)$ when $\mu < \mu_{k+2}$. Then again by Corollary 1,

$$\begin{aligned} u_{k+1}(\mu) &> u_{k+2}(\mu) \geq \mu(\tilde{g}_w + \delta\beta_w u_{k+1}^w(\phi^{-1}(\mu))) + (1-\mu)u_{k+2}^s \\ &= \mu\tilde{g}_w + (1-\mu)\tilde{g}_s + \delta(\mu\beta_w + (1-\mu)\beta_s)\mu_{k+1}(\phi^{-1}(\mu)) \\ &> \mu\tilde{g}_w + (1-\mu)\tilde{g}_s + \delta(\mu\beta_w + (1-\mu)\beta_s)h(\phi^{-1}(\mu)) \\ &= h(\mu). \end{aligned}$$

We can then conclude that (7) holds for all $\mu \in [\mu_{k+1}, \mu_{k+2})$ by the induction on j . And the proposition in the case $\beta_s > \beta_w$ follows by the induction on n .

If $\beta_s = \beta_w = \beta$, $\phi^{-1}(\mu) = \mu$. So if $\mu_n \leq \mu < \mu_{n+1}$, then by Lemma 3

$$\begin{aligned} u_n(\mu) &> u_{n+1}(\mu) \geq \mu(\tilde{g}_w + \delta\beta u_n^w(\phi^{-1}(\mu))) + (1-\mu)u_{n+1}^s \\ &= \mu(\tilde{g}_w + \delta\beta u_n^w(\mu)) + (1-\mu)u_{n+1}^s \\ &= \mu\tilde{g}_w + (1-\mu)\tilde{g}_s + \delta\beta u_n(\mu), \end{aligned}$$

which implies

$$u_n(\mu) > \frac{\mu\tilde{g}_w + (1-\mu)\tilde{g}_s}{1-\delta\beta} = h(\mu).$$

□

We shall show that PBE is generically unique in that the equilibrium paths are the same: given an initial belief $\mu \in (\mu_n, \mu_{n+1})$, G offers $x_n, x_{n-1}, \dots, x_0 = \underline{x}_s$ in turn along the path, R_w mixes between rejection and acceptance such that the belief path is $(\mu, \mu_{n-1}, \dots, \mu_0 = 0)$, and R_s always rejects until G offers $\underline{x}_s$.

For the following, we fix an arbitrary PBE of the game.

Lemma 5. *At any history, the maximum offer in the support of G 's strategy is no more than $\underline{x}_s$.*

Proof. Let $\bar{x}(h)$ denote the maximum offer in the support of G 's strategy at a history h , and denote by $\bar{x}$ the supremum of the union of G 's strategies over all histories, i.e., $\bar{x} = \sup_h \bar{x}(h)$. $\bar{x}$ must be finite, as G certainly would not offer anything larger than π . By way of contradiction, suppose $\bar{x} > \underline{x}_s$. So there exists a history h^t at which $\bar{x}(h^t) > \bar{x} - \epsilon$, where $\epsilon \triangleq (1 - \delta\beta_s)(\bar{x} - \underline{x}_s) > 0$. Note that if R does not accept an offer now, its payoff will be upper bounded by

$$\begin{aligned} (1-\delta)r_s + \delta(\alpha_s\pi + \beta_s\bar{x}) &= (1-\delta)r_s + \delta\alpha_s\pi + \delta\beta_s\bar{x} \\ &= (1-\delta\beta_s)\underline{x}_s + \delta\beta_s\bar{x} \\ &= (1-\delta\beta_s)(\underline{x}_s - \bar{x}) + \bar{x} \\ &= \bar{x} - \epsilon. \end{aligned}$$

As a result, any offer $x^t \in (\bar{x} - \epsilon, \bar{x}]$ will be accepted by R at time t . But then $\bar{x}(h^t)$ should not be larger than $\bar{x} - \epsilon$ as no offer in $(\bar{x} - \epsilon, \bar{x}]$ is a best response by G , a contradiction. It follows that we must have $\bar{x} \leq \underline{x}_s$. $\square$

Corollary 2. *If the current belief is $\mu^t \in [0, \mu_1)$, then G offers $\underline{x}_s$, and R accepts it.*

Proof. Immediate from the preceding lemma. $\square$

Lemma 6.

- (1) *Suppose the current belief is $\mu^t \in (\mu_1, \mu_2)$. Then G 's best response is x_1 . R_w accepts the offer while R_s rejects it.*
- (2) *Suppose the current belief is $\mu^t = \mu_1$. Then G 's best response is any mixture between x_1 and $\underline{x}_s$. If G offers x_1 , R_w accepts it with probability 1, and $\underline{x}_s$ rejects it. If G offers $\underline{x}_s$, R accepts it.*

Proof. The proof works slightly differently in two cases: (a) $\beta_s > \beta_w$, and (b) $\beta_s = \beta_w = \beta$. We provide the full proof for case (a), and give the mechanics and differences for case (b) when differences arise.⁴⁰

Suppose $\beta_s > \beta_w$ (case a), which implies that $\mu > \psi^{-1}(\mu)$ for all $\mu \in (0, 1)$. Let $\epsilon > 0$ be such that $\mu - \psi^{-1}(\mu) > \epsilon$ for all $\mu \in [\mu_1, \mu_2]$. Without loss of generality, we assume $N_1 \triangleq (\mu_2 - \mu_1)/\epsilon$ is an integer. First consider the case $\mu^t \in (\mu_1, \mu_1 + \epsilon)$. By assumption, $\mu^{t+1} < \mu_1$, so the next period's offer must be $\underline{x}_s$.

- i) First of all, note that G can guarantee a payoff arbitrarily close to $u_1(\mu^t)$ by making an offer slightly larger than x_1 in the current period (R_w will accept it for sure, since it will definitely be offered $\underline{x}_s$ in the next period).
- ii) By Lemma 5, G does not offer $x^t > \underline{x}_s$. And offering $x^t = \underline{x}_s$ gives payoff $u_0(\mu^t)$ which is strictly less than $u_1(\mu^t)$, and thus cannot be optimal.
- iii) The optimal offer x^t cannot be in $(x_1, \underline{x}_s)$, since G can lower x^t slightly while still keeping it larger than x_1 and gain a strictly higher payoff (since R_w accepts this any PBE).
- iv) The payoff for G from offering $x < x_1$ is

$$\mu^t (\tilde{g}_w + \delta\beta_w(\pi - \underline{x}_s)) + (1 - \mu^t)u_1^s < u_1(\mu^t),$$

by Lemma 2, Corollary 1. Thus $x^t < x_1$ cannot be optimal either.

- v) The only possible choice left is $x^t = x_1$. And R_w must accept it with probability 1, since if the rejection probability ν^t of R_w is strictly positive, then the payoff for G is

$$\mu^t ((1 - \nu^t)(\pi - x_1) + \nu^t(\tilde{g}_w + \delta\beta_w(\pi - \underline{x}_s))) + (1 - \mu^t)u_1^s,$$

which is strictly less than $u_1(\mu^t)$ by Lemma 2.

⁴⁰This Lemma, the next Lemma, and their proofs use the approach – “upwards induction on beliefs” – suggested in the sketch offered by Fudenberg and Tirole (1991, 409-410) for the two-type, infinite-horizon buyer-seller bargaining game. That game is close to being a special case of our model when $\beta_s = \beta_w = \beta$. In our main case of $\beta_s > \beta_w$, the analysis is actually simplified because the “mechanical” effect causes a belief change following rejection of an offer independent of R_w 's strategy.

In conclusion, if $\mu^t \in (\mu_1, \mu_1 + \epsilon)$, in equilibrium G offers $x^t = x_1$, R_w accepts it with probability 1, and R_s rejects it. If $\mu^t = \mu_1$, the only difference from above is step (ii), in which $u_0(\mu^t) = u_1(\mu^t)$. Thus in this case offering $\underline{x}_s$ is also optimal.

Now suppose that if $\mu^t \in (\mu_1, \mu_1 + k\epsilon)$, $1 \leq k < N_1$, in equilibrium G offers $x^t = x_1$ and R_w accepts it with probability 1. Consider the case $\mu^t \in [\mu_1 + k\epsilon, \mu_1 + (k+1)\epsilon)$. By assumption, μ^{t+1} is in either $[\mu_1, \mu_1 + k\epsilon)$ or $[0, \mu_1)$. If $\mu^{t+1} \in [0, \mu_1)$, then by the same reasoning as in (i) - (iii) above, the only possibility is that G offers x_1 , R_w accepts it with probability 1 and R_s rejects it. It remains to show that μ^{t+1} cannot be in $[\mu_1, \mu_1 + k\epsilon)$.

First assume $\mu^{t+1} \in (\mu_1, \mu_1 + k\epsilon)$. Again, by making an offer slightly larger than x_1 , G can get a payoff arbitrarily close to $u_1(\mu^t)$ (since R_w cannot get more by waiting). By the induction hypothesis, $x^{t+1} = x_1$, so x^t can only be x_2 , as otherwise R_w either accepts x^t for sure, which contradicts with the range of μ^{t+1} , or rejects x^t for sure, which makes G 's payoff less than $u_1(\mu^t)$. The payoff for G in equilibrium when facing R_w is then

$$(1 - \nu^t)(\pi - x_2) + \nu^t(\tilde{g}_w + \delta\beta_w(\pi - x_1)),$$

where $\nu^t = o(\mu^{t+1})\beta_s/o(\mu^t)\beta_w$. Note that ν^t is increasing in μ^{t+1} and

$$\pi - x_2 > \tilde{g}_w + \delta\beta_w(\pi - x_1).$$

Hence the payoff is strictly less than what it would be when $\mu^{t+1} = \mu_1$, which is exactly $u_2(\mu^t)$. However, $u_2(\mu^t)$ is less than G 's payoff guarantee $u_1(\mu^t)$ by definition of the recursion, since $\mu^t < \mu_2$ by hypothesis. This is impossible in equilibrium. It follows that we cannot have $\mu^{t+1} \in (\mu_1, \mu_1 + k\epsilon)$.

On the other hand, if $\mu^{t+1} = \mu_1$, then G 's payoff is either (if G offers x_2 at t) $u_2(\mu^t)$ or (if G offers x_1 at t)

$$(9) \quad \mu^t((1 - \nu_2(\mu^t))(\pi - x_1) + \nu_2(\mu^t)(\tilde{g}_w + \delta\beta_w(\pi - \underline{x}_s))) + (1 - \mu^t)u_1^s.$$

$u_2(\mu^t) < u_1(\mu^t)$, and (9) is strictly less than $u_1(\mu^t)$ by Lemma (2). So μ^{t+1} should not be μ_1 .

By induction on k , we have shown that, if the current belief is $\mu^t \in (\mu_1, \mu_2)$, then G offers x_1 , R_w accepts it, and $\underline{x}_s$ rejects it.

The second part of the lemma is clear by the definition of μ_1 . This proves the claim for case (a), $\beta_s > \beta_w$. We now sketch what needs to change in the case of $\beta_s = \beta_w = \beta$.

By Lemma 4, we can pick $\epsilon > 0$ such that

$$\frac{\mu(1 - \nu_\epsilon)\pi + \mu\nu_\epsilon\tilde{g}_w + (1 - \mu)\tilde{g}_s}{1 - \delta\beta} < u_1(\mu)$$

$$\mu(1 - \nu_\epsilon)\pi + \mu\nu_\epsilon\tilde{g}_w + (1 - \mu)\tilde{g}_s + \delta\beta u_1(\mu) < u_1(\mu),$$

for all $\mu \in (\mu_1, \mu_2)$, where $\nu_\epsilon = \inf_{\mu \in (\mu_1, \mu_2)} \{o(\mu - \epsilon)/o(\mu)\}$.⁴¹

First consider the case $\mu^t \in (\mu_1, \mu_1 + \epsilon)$.

⁴¹Note that $o(\mu - \epsilon)/o(\mu) = \frac{\frac{\mu - \epsilon}{1 - (\mu - \epsilon)}}{\frac{\mu}{1 - \mu}} = \left(1 - \frac{\epsilon}{\mu}\right) \frac{1 - \mu}{1 - \mu + \epsilon}$. So $\lim_{\epsilon \downarrow 0} \nu_\epsilon = 1$

- i) By Lemma 5, G will never make an offer larger than $\underline{x}_s$, and thus G can guarantee a payoff arbitrarily close to $u_1(\mu^t)$ by making an offer slightly larger than x_1 in the current period (because R_w would definitely accept).
- ii) Let $u^{\sup}$ denote the supremum of the expected payoffs over all equilibria with initial belief $\mu \in [\mu_1, \mu_1 + \epsilon)$. $u^{\sup}$ must be upper bounded by

$$\max \{u_1(\mu), \mu(1 - \nu_\epsilon)\pi + \mu\nu_\epsilon\tilde{g}_w + (1 - \mu)\tilde{g}_s + \delta\beta u^{\sup}\}.$$

(The first term is the maximum possible payoff when $\mu^{t+1} < \mu_1$ and the second term is a bound on the maximum possible payoff when $\mu^{t+1} \geq \mu_1$.)

The above maximum can only occur at the former expression, as otherwise

$$u^{\sup} \leq \frac{\mu(1 - \nu_\epsilon)\pi + \mu\nu_\epsilon\tilde{g}_w + (1 - \mu)\tilde{g}_s}{1 - \delta\beta} < u_1(\mu),$$

a contradiction. It follows that $u^{\sup} \leq u_1(\mu)$.

Suppose there exists an equilibrium in which $\mu^{t+1} \geq \mu_1$ with positive probability. Along any path on which $\mu^{t+1} \geq \mu_1$, G 's expected payoff from the current period onwards is upper bounded by

$$\begin{aligned} & \mu^t(1 - \nu_\epsilon)\pi + \mu^t\nu_\epsilon\tilde{g}_w + (1 - \mu^t)\tilde{g}_s + \delta\beta u^{\sup} \\ & \leq \mu^t(1 - \nu_\epsilon)\pi + \mu^t\nu_\epsilon\tilde{g}_w + (1 - \mu^t)\tilde{g}_s + \delta\beta u_1(\mu^t) \\ & < u_1(\mu^t), \end{aligned}$$

which is not optimal. And thus in any equilibrium μ^{t+1} must be less than μ_1 .

- iii) Other cases are the same as when $\beta_s > \beta_w$.

Upward induction from $(\mu_1, \mu_1 + k\epsilon)$ to $[\mu_1 + k\epsilon, \mu_1 + (k + 1)\epsilon)$ is likewise similar to the case of $\beta_s > \beta_w$.

□

Lemma 7. Suppose $n \geq 1$.

- (1) If the current belief is $\mu^t \in (\mu_n, \mu_{n+1})$, then G 's best response is offering x_n . If G offers x_n , R_w mixes between rejection and acceptance so that $\mu^{t+1} = \mu_{n-1}$, and R_s rejects it.
- (2) If the current belief is $\mu^t = \mu_n$, then G 's best response is any mixture between x_n and x_{n-1} . If G offers x_n , R_w mixes between rejection and acceptance so that $\mu^{t+1} = \mu_{n-1}$, and R_s always rejects. If G offers x_{n-1} and $n > 1$, R_w mixes between rejection and acceptance so that $\mu^{t+1} = \mu_{n-2}$. If $n = 1$ and G offers $\underline{x}_s$, R accepts the offer.

Proof. The argument is an extension to that for Lemma 6. We proceed by induction on n . The base case is just Lemma 6. Assuming the conclusion is true for all $n \leq j - 1$, where $j \geq 2$, we consider $\mu^t \in [\mu_j, \mu_{j+1})$. Again, we begin assuming that $\beta_s > \beta_w$, and then discuss modifications needed for the case of $\beta_s = \beta_w$ below.

Suppose $\beta_s > \beta_w$. Let $\epsilon > 0$ be such that $\mu^t - \psi^{-1}(\mu^t) > \epsilon$, and assume $N_j \triangleq (\mu_{j+1} - \mu_j)/\epsilon$ is an integer.

As a first step, consider the case $\mu^t \in (\mu_j, \mu_j + \epsilon)$. By assumption, $\mu^{t+1} < \mu_j$.

- i) We first note that G can guarantee itself a payoff arbitrarily close to $u_j(\mu^t)$ by offering $x^t = x_j + \eta$ for small $\eta > 0$. This is because by offering $x^t = x_j + \eta$, μ^{t+1} can only be μ_{j-1} ,⁴² which implies that G 's payoff is $u_j(\mu^t) - \mu^t(1 - \nu_n(\mu^t))\eta$, i.e., G 's continuation payoff remains the same as when $x^t = x_j$ but G pays an extra η in the event that R accepts.
- ii) Suppose $x^t = x_m$, where $m < j$, is a best response. An equilibrium entails $\mathbb{E}[x^{t+1}] = x_{m-1}$, as otherwise R_w would either accept for sure or reject for sure in this period, resulting in a payoff less than $u_j(\mu^t)$ for G . Since $\mu^{t+1} < \mu_j$, by induction we must have $\mu^{t+1} \in [\mu_{m-1}, \mu_m)$ in order to have $\mathbb{E}[x^{t+1}] = x_{m-1}$. However, by Lemma (2), if $\mu^{t+1} \in [\mu_{m-1}, \mu_m)$, G 's payoff would be upper bounded by $u_m(\mu^t)$, which is strictly less than $u_j(\mu^t)$, a contradiction.
- iii) Suppose $x^t \in (x_m, x_{m-1})$, where $m \leq j$, is a best response. By the same reason as for (i), we must have $\mu^{t+1} = \mu_{m-1}$. Hence G can strictly reduce payment by offering slightly less than x^t while providing the same incentive, which cannot happen in equilibrium.
- iv) An offer smaller than x_j will be rejected for sure, because the worst offer in the future is x_{j-1} . By Corollary (1), the payoff for G from offering $x < x_j$ is strictly less than $w_j(\mu^t)$, a contradiction.
- v) The only possibility left is $x^t = x_j$. To support equilibrium, μ^{t+1} can only be in $[\mu_{j-1}, \mu_j)$. If $\mu^{t+1} \in (\mu_{j-1}, \mu_j)$, by Lemma (2), the payoff is strictly less than $u_j(\mu^t)$. So we must have $\mu^{t+1} = \mu_{j-1}$.

If $\mu^t = \mu_j$, the only thing different from the case $\mu^t \in (\mu_j, \mu_j + \epsilon)$ is step (ii), since $u_m(\mu^t)$ is equal to $u_j(\mu^t)$ if $m = j - 1$. And thus G 's best response could be a mixture between x_j and x_{j-1} . If G offers x_j , the situation is the same as above. If G offers x_{j-1} , the only possible case in which G still gets $u_j(\mu^t)$ is when $\mu^{t+1} = \mu_{j-2}$.

Next suppose we have shown that if $\mu^t \in (\mu_j, \mu_j + (k - 1)\epsilon)$, $2 \leq k \leq N_j$, then G 's best response is x_n , and $\mu^{t+1} = \mu_{j-1}$. Then consider the case $\mu^t \in [\mu_j + (k - 1)\epsilon, \mu_j + k\epsilon)$. Again, by definition of ϵ , $\mu^{t+1} < \mu_j + (k - 1)\epsilon$. If $\mu^{t+1} < \mu_j$, by the same reasoning as in (i) - (v) above, the only possibility is that G offers $x^t = x_j$, and $\mu^{t+1} = \mu_{j-1}$. Moreover, by the argument in (i), G can guarantee itself a payoff arbitrarily close to $u_j(\mu^t)$. It remains to show that μ^{t+1} cannot be larger than or equal to μ_j .

Now assume $\mu^{t+1} > \mu_j$. By the induction hypothesis, $x^{t+1} = x_j$, so x^t can only be x_{j+1} . The payoff for G when facing R_w is

$$(1 - \nu^t)(\pi - x_{j+1}) + \nu^t(\tilde{g}_w + \delta\beta_w u_j^w(\mu^{t+1})),$$

which, by Lemma (2), is at most $u_{j+1}^w(\mu^t)$. By definition of μ_{j+1} , $u_{j+1}^w(\mu^t) < u_j^w(\mu^t)$. And thus if $\mu^{t+1} > \mu_j$, the expected total payoff from period t onwards is strictly less than $u_j^w(\mu^t)$, which is absurd.

⁴²If $\mu_{j-1} < \mu^{t+1} < \mu_j$, by the induction hypothesis, the next period's offer must be x_{j-1} , accepting which is strictly worse than accepting $x_j + \eta$ in the current period for R_w . If R_w accepts $x_j + \eta$ for sure, $\mu^{t+1} = 0$. If $\mu^{t+1} < \mu_{j-1}$, the next period's offer is at least x_{j-2} , accepting which is strictly better than accepting $x_j + \eta$, but then if R_w rejects $x_j + \eta$ for sure, $\mu^{t+1} = \psi^{-1}(\mu^t) > \mu_{j-1}$.

If $\mu^{t+1} = \mu_j$, the payoff is either (if G offers x_{j+1} at t) $u_{j+1}(\mu^t)$ or (if G offers x_j at t)

$$(10) \quad \mu^t \left[((1 - \nu_{j+1}(\mu^t))(\pi - x_j) + \nu_{j+1}(\mu^t)(\tilde{g}_w + \delta\beta_w u_{j-1}^w(\mu_j))) \right] + (1 - \mu^t)u_j^s.$$

$u_{j+1}(\mu^t) < u_j(\mu^t)$, and (10) is strictly less than $u_j(\mu^t)$ by Lemma (2). So μ^{t+1} cannot be μ_j .

By induction on k , the desired conclusion holds for $n = j$. And the lemma then follows by induction on j .

Last, we consider the case of $\beta_s = \beta_w = \beta$. Again, the method is induction: assuming the conclusion is true for all $n \leq j - 1$, where $j \geq 2$, we consider $\mu^t \in [\mu_j, \mu_{j+1})$.

The additional difficulty over the case of $\beta_s > \beta_w$ is that we need a finer bound on what G can offer when $\mu^t \in [\mu_j, \mu_{j+1})$. So to start, we shall show that in any equilibrium and in any period t' following any history, $\mathbb{E}[x^{t'}] \leq x_j$ when $\mu^{t'} \geq \mu_j$. Let $x^{\sup}$ denote the supremum of the expected offer G will propose when the current belief is larger than μ_j over all equilibria. $x^{\sup}$ is finite as it must be upper bounded by $\underline{x}_s$ due to Lemma 5. Suppose $x_k < x^{\sup} \leq x_{k-1}$ for some $k \leq j$. For any $\epsilon_2 > 0$ there exists some equilibrium in which following some history $\mathbb{E}[x^\tau] > x^{\sup} - \epsilon_2$ and $\mu^\tau \geq \mu_j$. We pick an $\epsilon_2 > 0$ such that $x^{\sup} - \epsilon_2 > x_k$, and will show that any offer $x^\tau > x^{\sup} - \epsilon_2$ is strictly worse than $x^\tau = x^{\sup} - \epsilon_2$ for G . First note that proposing an offer larger than or equal to $x^{\sup} - \epsilon_2$ must lead to a belief $\mu^{\tau+1} < \mu_j$ as R_w definitely prefers the current offer than an expected offer of at most $x^{\sup}$ in the next period. That $\mu^{\tau+1} < \mu^\tau$ implies R_w accepts such an offer with positive probability. There are two cases to consider.

- If $k = 1$, for R_w any $x^\tau \geq x^{\sup} - \epsilon_2$ in this period is strictly better than $\underline{x}_s$ in next period, and thus R_w must accept such x^τ almost surely, leading to $\mu^{\tau+1} = 0$. In this case, any offer $x^\tau > x^{\sup} - \epsilon_2$ is clearly strictly worse than $x^\tau = x^{\sup} - \epsilon_2$ for G .
- Suppose $k > 1$ instead. In any equilibrium, if $x^\tau = x^{\sup} - \epsilon_2$, $\mu^{\tau+1}$ must not be 0 as $\underline{x}_s$ in next period is strictly better for R_w . That $0 < \mu^{\tau+1} < \mu^\tau$ implies R_w must be indifferent between x^τ in this period and the expected offer in next period, i.e., $x^\tau = (1 - \delta)r_w + \delta(\alpha_w\pi + \beta_w\mathbb{E}[x^{\tau+1}])$. The only belief at which G can offer such $\mathbb{E}[x^{\tau+1}]$ is $\mu^{\tau+1} = \mu_{k-1}$. Thus the expected payoff for G by offering $x^\tau = x^{\sup} - \epsilon_2$ is

$$(11) \quad \mu^\tau \left((1 - \nu_k(\mu^\tau))(\pi - (x^{\sup} - \epsilon_2)) + \nu_k(\mu^\tau)(\tilde{g}_w + \delta\beta_w u_{k-1}^w(\mu_{k-1})) \right) + (1 - \mu^\tau)u_k^s.$$

If $x^\tau > x^{\sup} - \epsilon_2$, $\mu^{\tau+1}$ cannot be 0 either. So we must have $x^\tau = (1 - \delta)r_w + \delta(\alpha_w\pi + \beta_w\mathbb{E}[x^{\tau+1}])$ as well. This implies $\mu^{\tau+1} \leq \mu_{k-1}$. If $\mu^{\tau+1} = \mu_{k-1}$, G 's expected payoff is

$$\mu^\tau \left((1 - \nu_k(\mu^\tau))(\pi - x^\tau) + \nu_k(\mu^\tau)(\tilde{g}_w + \delta\beta_w u_{k-1}^w(\mu_{k-1})) \right) + (1 - \mu^\tau)u_k^s < (11).$$

And if $\mu_{\ell-1} \leq \mu^{\tau+1} < \mu_\ell$ for some $\ell \leq k - 1$, G 's expected payoff is

$$(12) \quad \mu^\tau \left(\left(1 - \frac{o(\mu^{\tau+1})}{o(\mu^\tau)} \right) (\pi - x^\tau) + \frac{o(\mu^{\tau+1})}{o(\mu^\tau)} (\tilde{g}_w + \delta\beta_w u_{\ell-1}^w(\mu_{\ell-1})) \right) + (1 - \mu^\tau)u_\ell^s$$

$$(13) \quad \leq \mu^\tau \left((1 - \nu_\ell(\mu^\tau))(\pi - x^\tau) + \nu_\ell(\mu^\tau)(\tilde{g}_w + \delta\beta_w u_{\ell-1}^w(\mu_{\ell-1})) \right) + (1 - \mu^\tau)u_\ell^s$$

$$< (11)$$

where (12) is due to [INSERT: Lemma 2] and (13) is due to Lemma 3.

In either case, $x^\tau > x^{\sup} - \epsilon_2$ is strictly worse than $x^\tau = x^{\sup} - \epsilon_2$ for G . So in equilibrium we must have $Pr((x^\tau > x^{\sup} - \epsilon_2) = 0)$, which contradicts that $\mathbb{E}[x^\tau] > x^{\sup} - \epsilon_2$. It then follows that $\mathbb{E}[x^{t'}] \leq x_j$ in any equilibrium when $\mu^{t'} \geq \mu_j$. In particular, we have $\mathbb{E}[x^t] \leq x_j$.

Next we turn to upward induction to characterize the equilibrium strategies of G . By Lemma 4, we can pick $\epsilon > 0$ such that

$$\frac{\mu(1 - \nu_\epsilon)\pi + \mu\nu_\epsilon\tilde{g}_w + (1 - \mu)\tilde{g}_s}{1 - \delta\beta} < u_j(\mu)$$

$$\mu(1 - \nu_\epsilon)\pi + \mu\nu_\epsilon\tilde{g}_w + (1 - \mu)\tilde{g}_s + \delta\beta u_j(\mu) < u_j(\mu),$$

for all $\mu \in (\mu_j, \mu_{j+1})$, where $\nu_\epsilon = \inf_{\mu \in (\mu_j, \mu_{j+1})} \{o(\mu - \epsilon)/o(\mu)\}$.

First consider the case $\mu^t \in (\mu_j, \mu_j + \epsilon)$.

- (1) We have shown that the expected offer in any period cannot exceed x_j , when the belief in that period is larger than or equal to μ_j . As a result, if $x^t > x_j$, μ^{t+1} must be less than μ_j . Hence G can guarantee itself a payoff arbitrarily close to $u_j(\mu^t)$ by offering $x^t = x_j + \eta$ for small $\eta > 0$ for the same reason given in the proof of Lemma 5.
- (2) Other parts are similar to the $\beta_s > \beta_w$ case.

□

Lemma 8 (Uniqueness of Equilibrium). *Suppose the initial belief is $\mu^0 \in [\mu_n, \mu_{n+1})$ for some $n \geq 1$. Then the equilibrium path of any PBE must satisfy: At $t = 0$, G offers x_n if $\mu^0 > \mu_n$, and G mixes between x_n and x_{n-1} if $\mu^0 = \mu_n$.*

- (1) If G offers x_n at $t = 0$, then the bargaining lasts at most $n + 1$ rounds. At each $t = 1, 2, \dots, n$, the belief at the beginning of the period is $\mu^t = \mu_{n-t}$, and G 's offer is x_{n-t} .
- (2) If G offers x_{n-1} at $t = 0$, then the bargaining lasts at most n rounds. At each $t = 1, 2, \dots, n-1$, the belief at the beginning of the period is $\mu^t = \mu_{n-t-1}$, and G 's offer is x_{n-t-1} .

Proof. The result is an immediate consequence of Lemma 6 and Lemma 7. □

Step 3: Example of a PBE.

To show the existence, we describe explicitly a PBE as follows.

- Suppose there was no deviation in history. If $\mu^t \in [\mu_n, \mu_{n+1})$, then G offers $x^t = x_n$. If $n \geq 1$, R_w mixes between rejection and acceptance such that $\mu^{t+1} = \mu_{n-1}$, and R_s rejects the offer.⁴³ If $n = 0$, both R_w and R_s accept the offer.
- Suppose there has been deviation in history, the belief is $\mu^t \in [\mu_n, \mu_{n+1})$, and the last offer (x^t for R while x^{t-1} for G) is x satisfying $x_{k+1} \leq x < x_k$ for some $k \geq 0$.

R's turn: i) If $k > n$, R rejects the offer.
 ii) If $k = n$ and $\mu^t < \psi(\mu_n)$, R rejects the offer.
 iii) If $k = n$ and $\mu^t \geq \psi(\mu_n)$, R_w mixes between rejection and acceptance such that $\mu^{t+1} = \mu_n$. R_s rejects the offer.
 iv) If $1 \leq k < n$, R_w mixes between rejection and acceptance such that $\mu^{t+1} = \mu_k$. R_s rejects the offer.
 v) If $k = 0$, R_w accepts the offer while R_s rejects the offer.

G's turn: i) If $\mu^t > \mu_n$ or $n \neq k$, then G offers x_n .
 ii) If $\mu^t = \mu_n$ and $n = k \geq 1$, then G mixes between offering x_k and offering x_{k-1} with probabilities $(1 - \lambda, \lambda)$ such that

$$x = r_w + \alpha_w \pi + (1 - \lambda)\delta\beta_w x_k + \lambda\delta\beta_w x_{k-1}.$$
 iii) If $\mu^t = 0$, then G offers $\underline{x}_s$.
- Suppose there has been deviation in history and the last offer is $x \geq \underline{x}_s$.

R's turn: R accepts the offer.

G's turn: G offers x_n .

Step 4: Verification of the PBE.

Finally we verify that the above is indeed a PBE. G's decision is optimal by Lemma (7), and R_s 's decision is obvious. So we only need to verify the decisions of R_w .

- Suppose there was no deviation in history, $\mu^t \in [\mu_n, \mu_{n+1})$, and $x^t = x_n$ is offered by G. By construction of x_k 's, R_w is always indifferent between x_n and any other offer in the future, so any strategy of R_w is optimal.
- Suppose there has been deviation in history, and $\mu^t \in [\mu_n, \mu_{n+1})$. We consider each case (i) - (v) specified in the strategies.
 - The current offer is $x < x_{n+1}$, which is strictly worse than accepting the offer (either x_n or x_{n-1}) in next period.
 - The current offer is $x < x_n$, which is strictly worse than accepting x_{n-1} in next period.
 - Upon rejection, next period G will mix between offering x_n and offering x_{n-1} such that the expected payoff of accepting the offer next period, by construction, is equal to the payoff of accepting x^t this period. So any strategy of R_w is optimal.
 - Same as in (iii), next period G will mix in a way such that R_w is indifferent between acceptance this period and acceptance next period. So any strategy of R_w is optimal.
 - The best offer in the future is $\underline{x}_s$, which is no better than accepting x^t this period.

⁴³Note that at $n = 1$ the mix is degenerate in that $\nu_1 = 0$.

- Suppose there was deviation in history and the last offer is $x \geq \underline{x}_s$. Clearly acceptance is optimal for R.

2. PROPOSITION 2

Proposition 2. *Let Δ be the time between offers in the armed conflict bargaining game when $\lambda_s \leq \lambda_w$, and let $N(\Delta)$ be the maximum number of rounds of bargaining in equilibrium. Then the total possible maximum duration of war, $N(\Delta)\Delta$, approaches zero as Δ approaches zero if and only if either (a) $\pi - w_w \geq \underline{x}_s$, which is implied by $A \geq 1$, or (b) $\pi - w_w < \underline{x}_s$ and $\mu < \mu^*$ defined by*

$$o^* = \frac{\mu^*}{1 - \mu^*} = \frac{o_1}{1 - A} = \frac{\rho + \lambda_s}{\rho + \lambda_w} \frac{\pi - w_s - \underline{x}_s}{\underline{x}_s - (\pi - w_w)}.$$

If (c) $\pi - w_w < \underline{x}_s$ and $\mu > \mu^*$, then the maximum and mean war durations are bounded away from zero in the limit as $\Delta \rightarrow 0$. When $\lambda_s = \lambda_w$ (i.e., $\beta_s = \beta_w$), limiting maximum duration is $2\hat{T}$, where $\hat{T} = \log(1/A)/(\rho + \lambda_w)$. When $\lambda_s > \lambda_w$, limiting maximum duration approaches $\hat{T}$ from above.

Proof of Proposition 2

Preliminaries.

- Let $r \triangleq \beta_s/\beta_w = e^{\Delta(\lambda_w - \lambda_s)}$. Note $r \geq 1$.
- Let $L \triangleq \lambda_w - \lambda_s$.
- Let $\rho' \triangleq \rho + \lambda_w$.
- Define the sequence $O_n(x)$ for $n > 1$, $x > 0$, and initial conditions $O_1 > O_0 \geq 0$ by

$$O_n = r(1 + x)O_{n-1} - rxO_{n-2}.$$

We note first that from

$$O_n - O_{n-1} = rx(O_{n-1} - O_{n-2}) + (r - 1)O_{n-1}$$

it is immediate that for all $n > 0$, $O_n > O_{n-1}$ and

$$(14) \quad O_n(x) > O_n(x') \text{ whenever } x > x'.$$

$O_n(x)$ is a second-order linear difference equation with roots of its characteristic equation

$$m_1(x) = \frac{1}{2} \left(r(1 + x) + \sqrt{r^2(1 + x)^2 - 4rx} \right) \text{ and } m_2(x) = \frac{1}{2} \left(r(1 + x) - \sqrt{r^2(1 + x)^2 - 4rx} \right).$$

Using the standard method, we have that its general solution is

$$(15) \quad O_n = \frac{1}{m_1 - m_2} [m_1^n(O_1 - O_0 m_2) - m_2^n(O_1 - O_0 m_1)]$$

or, in the case of $O_0 = 0$,

$$O_n = O_1 \frac{m_1^n - m_2^n}{m_1 - m_2}.$$

Lemma 9. Let $n = T/\Delta$.⁴⁴ All the following limits are for $\Delta \downarrow 0$.

- (1) $\lim m_1 = 1$ and $\lim m_2 = x$.
- (2) $\lim m_2^{T/\Delta} = 0$ if $x < 1$ and positive infinity if $x > 1$.
- (3) For $x \neq 1$, $\lim m_1^{T/\Delta} = \exp\left(\frac{TL}{|1-x|}\right)$.
- (4) Given $O_0 = 0$ and for any $T > 0$, if $x > 1$ then $\lim O_{\lfloor T/\Delta \rfloor} = +\infty$. If $x < 1$, then

$$\lim O_{\lfloor T/\Delta \rfloor} = \frac{O_1}{1-x} \exp\left(\frac{TL}{1-x}\right).$$

Proof. Items 1 and 2 are immediate. For 3, consider $T \log m_1(x)/\Delta$. Using L'Hôpital's rule,

$$\lim_{\Delta \downarrow 0} T \frac{\log m_1(x)}{\Delta} = T \lim_{\Delta \downarrow 0} \frac{\partial}{\partial \Delta} \log m_1(x) = T \lim_{\Delta \downarrow 0} \frac{dm_1(x)/d\Delta}{m_1(x)}.$$

Work shows that the final term on the right resolves to⁴⁵

$$\begin{aligned} T \lim \frac{dm_1(x)}{d\Delta} &= \frac{T}{2} \left[L(1+x) + \frac{2L(1+x)^2 - 4Lx}{2\sqrt{(1+x)^2 - 4x}} \right] \\ &= \frac{TL}{2} \left[\frac{(1+x)\sqrt{(1-x)^2 + 1 + x^2}}{\sqrt{(1-x)^2}} \right] \\ &= \frac{TL}{2} \left[\frac{1 - x^2 + 1 + x^2}{|1-x|} \right] = \frac{TL}{|1-x|}, \end{aligned}$$

from which (3) follows. Item 4 is implied by plugging into $\lim O_{\lfloor T/\Delta \rfloor}$ using the preceding results.

Lemma 10. If $A > 1$, then for any $\epsilon > 0$, $\lim o_{\lfloor \epsilon/\Delta \rfloor} = +\infty$. If $A < 1$, then for any $\epsilon > 0$ such that $Ae^{\rho'\epsilon} < 1$,

$$o^* < \frac{o_1}{1-A} \exp\left(\frac{\epsilon L}{1-A}\right) < \lim o_{\lfloor \epsilon/\Delta \rfloor} < \frac{o_1}{1-Ae^{\rho'\epsilon}} \exp\left(\frac{\epsilon L}{1-Ae^{\rho'\epsilon}}\right).$$

Proof. For $A > 1$, from (14), $o_{\lfloor \epsilon/\Delta \rfloor} > O_{\lfloor \epsilon/\Delta \rfloor}(A)$ since $A_n > A$ for all $n > 1$. By Lemma 9(4), $\lim O_{\lfloor \epsilon/\Delta \rfloor}$ is positive infinity, implying that this is true for $o_{\lfloor \epsilon/\Delta \rfloor}$ also. For $A < 1$, we have that, for n such that $0 < n\Delta < \epsilon$ and $O_1 = o_1$, $O_0 = 0$,

$$O_{\lfloor \epsilon/\Delta \rfloor}(A) < o_{\lfloor \epsilon/\Delta \rfloor} < O_{\lfloor \epsilon/\Delta \rfloor}(Ae^{\rho'\epsilon}),$$

⁴⁴We can ignore integer issues in what follows.

⁴⁵The final expression does not depend on the limit value of $dx/d\Delta$. Since below we use $x = A(\Delta)$, we need to show the existence of the limit. Through repeated application of L'Hôpital's rule and lengthy algebra, we can show

$$\lim_{\Delta \downarrow 0} \frac{dA}{d\Delta} = -\frac{(A^*)^2}{2\bar{v}} \left(\frac{\lambda_w + \rho}{\lambda_s + \rho} (\lambda_w - \lambda_s) r_s + \frac{\lambda_w + \rho}{\lambda_s + \rho} \frac{\lambda_w + \rho}{\rho} \lambda_s \tilde{\alpha}_s \pi - \frac{\lambda_w + \rho}{\rho} \lambda_w \tilde{\alpha}_w \pi \right).$$

Thus $\lim_{\Delta \downarrow 0} dA/d\Delta$ exists and is finite.

because $A < A_{\lfloor \epsilon/\Delta \rfloor}$ and $Ae^{\rho' n \Delta} < Ae^{\rho' \epsilon}$. Taking limits as $\Delta \downarrow 0$ and using Lemma 9(4) we have the conclusion.

Coase conjecture claims

We can now demonstrate the results concerning the Coase conjecture in the Proposition.

Case 1: $A > 1$. Assume to the contrary that there exists some initial belief $o > 0$ and a $T > 0$ such that $\lim o_{\lfloor T/\Delta \rfloor} < o$. (That is, assume to the contrary that maximum bargaining duration can be strictly positive in the limit as $\Delta \rightarrow 0$.) Since $A > 1$, using Lemma 10 we can always choose an ϵ such that $0 < \epsilon < T$ and have that $\lim O_{\lfloor \epsilon/\Delta \rfloor}(A) = +\infty$, which implies that there exists $\bar{\Delta}$ such that for all $\Delta < \bar{\Delta}$, $o_{\lfloor \epsilon/\Delta \rfloor} > o$, since $o_{\lfloor T/\Delta \rfloor} > O_{\lfloor T/\Delta \rfloor}(A)$. Thus there cannot be a strictly positive maximum duration in the limit as $\Delta \rightarrow 0$.

Case 2: $A < 1$ and $\mu < \mu^*$. From Lemma 10, for any $o < o^*$ and any small enough $\epsilon > 0$, $o < o_{\lfloor \epsilon/\Delta \rfloor}$ for all $\Delta < \bar{\Delta}$ for some $\bar{\Delta} > 0$. By the same kind of argument as in Case 1, it cannot be that maximum bargaining duration is strictly positive in the limit as $\Delta \downarrow 0$.

Case 3: $A < 1$ and $\mu > \mu^*$. For convenience let $\rho' \equiv \rho + \lambda_w$. Choose any $\mu \in (\mu^*, 1)$, thus $o > o^*$. Let T solve

$$o = \frac{o_1}{1 - Ae^{\rho' T}} \exp\left(\frac{TL}{1 - Ae^{\rho' T}}\right).$$

Note that $o > o^*$ implies $T > 0$, and $Ae^{\rho' T} < 1$. Let $B = Ae^{\rho' T}$. For any Δ and all $n < T/\Delta$, $O_n(B) > o_n$, since $B > Ae^{\rho' n \Delta}$ (using (14)). Because the term on the right-hand side above is the limit of $O_{\lfloor T/\Delta \rfloor}(B)$ as $\Delta \rightarrow 0$ (with $O_1 = o_1$ and $O_0 = 0$), and $O_{\lfloor T/\Delta \rfloor}(B) > o_{\lfloor T/\Delta \rfloor}$, it follows that the limiting maximum duration of bargaining is at least $n\Delta = T > 0$ for sufficiently small Δ . This proves that the Coase conjecture fails if $A < 1$ and $o > o^*$.

To see that mean duration also is bounded away from zero in this case, note that in the event of fighting against a strong type, the probability that the bargaining ends with a negotiated settlement (i.e., ends in period $N(\Delta)$) is $\beta_s^{N(\Delta)}$. Hence the mean duration is bounded below by

$$(1 - \mu)\beta_s^{N(\Delta)}(N(\Delta) + 1)\Delta = (1 - \mu) \exp(-\lambda_s N(\Delta)\Delta) (N(\Delta) + 1)\Delta,$$

which is strictly positive if $0 < \lim_{\Delta \downarrow 0} N(\Delta)\Delta < \infty$. But we just showed that for this case $0 < T \leq \lim N(\Delta)\Delta$, with T defined above. And for any initial belief o , we can choose a $T' > T$ such that $Ae^{\rho' T'} < 1$ and

$$o < \frac{o_1}{1 - Ae^{\rho' T'}} \exp\left(\frac{T' L}{1 - Ae^{\rho' T'}}\right).$$

Thus, for small enough Δ bargaining is certainly finished in finite time T' , proving that mean duration is strictly positive.

Limiting maximum durations

Next we prove the claims about limiting maximum durations, beginning with the case of $A < 1$ and $\lambda_w = \lambda_s$.

With $\lambda_w = \lambda_s$, we can express o_n as a manageable sum, as follows. The recursion in this case can be written $o_n - o_{n-1} = A_n(o_{n-1} - o_{n-2})$. Letting $z_n = o_n - o_{n-1}$,

$$z_n = \frac{A}{(\delta\beta_w)^{n-1}} z_{n-1} \text{ for } n > 1, \text{ where } z_1 = o_1, \text{ so that, for } n \geq 1$$

$$z_n = \frac{A^{n-1}}{(\delta\beta_w)^{\frac{n(n-1)}{2}}} o_1.$$

Using $\sum_{i=1}^n z_i = o_n$, we get that

$$o_n = o_1 \sum_{i=1}^n A^{i-1} (\delta\beta_w)^{-\frac{i(i-1)}{2}}.$$

Rewriting in terms of the time length of a period $\Delta > 0$ (and recalling $\rho' = \lambda_w + \rho$),

$$(16) \quad \begin{aligned} o_n &= o_1 \sum_{i=1}^n A^{i-1} e^{\rho' \Delta i(i-1)/2} \\ &= o_1 \sum_{i=1}^n \exp[(i-1) \log A + \rho' \Delta i(i-1)/2] \end{aligned}$$

Choose $\epsilon > 0$ small enough that $Ae^{\rho' \epsilon/2} < 1$. We first show that $o_1/(1-A) < \lim_{\Delta \downarrow 0} o_{\lfloor \epsilon/\Delta \rfloor} < o_1/(1-Ae^{\rho' \epsilon/2})$.

$$\begin{aligned} o_1 \sum_{i=1}^{\lfloor \epsilon/\Delta \rfloor} A^{i-1} &< o_{\lfloor \epsilon/\Delta \rfloor} = o_1 \sum_{i=1}^{\lfloor \epsilon/\Delta \rfloor} (Ae^{\rho' \Delta i/2})^{i-1} \\ o_1 \frac{1-A^{\epsilon/\Delta}}{1-A} &< o_{\lfloor \epsilon/\Delta \rfloor} < o_1 \sum_{i=1}^{\lfloor \epsilon/\Delta \rfloor} (Ae^{\rho' \epsilon/2})^{i-1} = o_1 \frac{1-(Ae^{\rho' \epsilon/2})^{\epsilon/\Delta}}{1-Ae^{\rho' \epsilon/2}}. \end{aligned}$$

Taking limits as $\Delta \downarrow 0$ establishes the claim.

Next, let $f(i)$ be the quadratic function $(i-1)(\log A + \rho' \Delta i/2)$, which is the exponent in (16). $f(i)$ has roots at $i = 1$ and $i = -2 \log A / \rho' \Delta = 2\hat{T}/\Delta$, and is negative for values in between these roots. Let $\bar{n} = (2\hat{T} - \epsilon)/\Delta$ and $\underline{n} = \epsilon/\Delta$ for $0 < \epsilon < 2\hat{T}$ and ϵ is chosen to produce integer values. Note that $f(\bar{n}) = f(\underline{n}) < 0$. So

$$\begin{aligned} S &\equiv \sum_{i=\underline{n}}^{\bar{n}} \exp(f(i)) < (\bar{n} - \underline{n} + 1) \exp(f(\underline{n})), \text{ and} \\ 0 &\leq \lim_{\Delta \downarrow 0} S < \lim_{\Delta \downarrow 0} \left(\frac{2(\hat{T} - \epsilon)}{\Delta} + 1 \right) \exp \left(\left(\frac{\epsilon}{\Delta} - 1 \right) (\log A + \frac{\rho' \epsilon}{2}) \right) = 0, \end{aligned}$$

so the limit of $S = 0$. (The last limit can be shown to be zero, coming from the fact that the power term goes to negative infinity.)

It follows from the above that we can choose $\epsilon > 0$ as close as we want to zero, and then for $\Delta \downarrow 0$, both $o_{\lfloor \epsilon/\Delta \rfloor}$ and $o_{\lfloor (2\hat{T}-\epsilon)/\Delta \rfloor}$ approach $o^* = o_1/(1-A)$. Thus o_n approaches a step function that goes immediately to o^* and is then flat up $2\hat{T}$.

Because $A_{\lfloor (2\hat{T}+\epsilon)/\Delta \rfloor} > 1$ for small enough Δ and any $\epsilon > 0$ and A_n is increasing, using the analysis above for $x > 1$, it follows that $o_{\lfloor (2\hat{T}+\epsilon)/\Delta \rfloor}$ goes to infinity for any $\epsilon > 0$. Thus maximum limiting duration in this case, for $o > o^* = o_1/(1-A)$, is $2\hat{T}$.

Next, the case of $\lambda_w > \lambda_s$. We will show that just after $\hat{n} \approx \hat{T}/\Delta$ (so that $A_n > 1$ for $n > \hat{n}$), $O_n(A_{\hat{n}})$ goes to infinity within any positive time duration ϵ for small enough Δ , and therefore o_n does also. Recall that for the series $O_n(x)$, if $O_1 > O_0 > 0$, the general solution is

$$O_n = \frac{1}{m_1 - m_2} [m_1^n (O_1 - O_0 m_2) - m_2^n (O_1 - O_0 m_1)].$$

Using Lemma 9, for any $x > 1$,

$$\lim_{\Delta \downarrow 0} O_{\lfloor T/\Delta \rfloor} = \frac{e^{TL/(x-1)}}{1-x} \lim (O_1 - O_0 x) - \frac{1}{1-x} \lim x^{T/\Delta} (O_1 - O_0).$$

We show that for $O_0 \equiv o_{\lfloor \hat{T}/\Delta \rfloor}$ and $O_1 \equiv o_{\lfloor \hat{T}/\Delta + 1 \rfloor}$, $\lim x^{T/\Delta} (O_1 - O_0)$ is positive infinity, which implies the claim.

From the recursion

$$o_{\lfloor \hat{T}/\Delta + 1 \rfloor} - o_{\lfloor \hat{T}/\Delta \rfloor} = r A_{\lfloor \hat{T}/\Delta \rfloor} (o_{\lfloor \hat{T}/\Delta \rfloor} - o_{\lfloor \hat{T}/\Delta - 1 \rfloor}) + (r-1) o_{\lfloor \hat{T}/\Delta \rfloor} > (r-1) o_{\lfloor \hat{T}/\Delta \rfloor} > (r-1) o^*,$$

so that

$$\lim x^{T/\Delta} (o_{\lfloor \hat{T}/\Delta + 1 \rfloor} - o_{\lfloor \hat{T}/\Delta \rfloor}) > \lim x^{T/\Delta} (e^{L\Delta} - 1) o^*.$$

With multiple applications of L'Hôpital's rule, the final expression is shown to be positive infinity for any $T > 0$ (when Δ approaches zero from above).

Finally, we derive the continuous approximation to o_n discussed in the text for the case of $\lambda_w > \lambda_s$.

Let $\tau \geq 0$ and $n = \lfloor \tau/\Delta \rfloor$. Then $o_{\lfloor \tau/\Delta \rfloor} = o_n$ is a right-continuous step function in τ , with upward jumps at points $\tau = \Delta, 2\Delta, 3\Delta, \dots$. The slope in τ of the line connecting points $(\lfloor \tau/\Delta \rfloor, o_{\lfloor \tau/\Delta \rfloor})$ and $(\lfloor \tau/\Delta + 1 \rfloor, o_{\lfloor \tau/\Delta + 1 \rfloor})$ is $s_n = (o_{n+1} - o_n)/\Delta$.

Rewriting the main recursion, using the approximation $s_{\lfloor \tau/\Delta \rfloor} \approx s_{\lfloor \tau/\Delta - 1 \rfloor}$ for small Δ , and dividing by Δ we have

$$\begin{aligned} o_{n+1} - o_n &= rA_n(o_n - o_{n-1}) + (r-1)o_n \\ o_{n+1} - o_n &\approx rA_n(o_{n+1} - o_n) + (r-1)o_n \\ o_{n+1} - o_n &= \frac{r-1}{1-rA_n}o_n \\ \frac{o_{n+1} - o_n}{\Delta} &= \frac{(r-1)/\Delta}{1-rA_n}o_n \\ \frac{o_{\lfloor \tau/\Delta + 1 \rfloor} - o_{\lfloor \tau/\Delta \rfloor}}{\Delta} &= \frac{(r-1)/\Delta}{1-rA_{\lfloor \tau/\Delta \rfloor}}o_{\lfloor \tau/\Delta \rfloor} \end{aligned}$$

Taking limits yields

$$o'(\tau) = \frac{L}{1 - Ae^{\rho'\tau}}o(\tau),$$

where $o(\tau)$ is a continuous function whose derivative at τ equals the limit of $s_{\lfloor \tau/\Delta \rfloor}$.

We can then solve for $o(\tau)$ as follows.

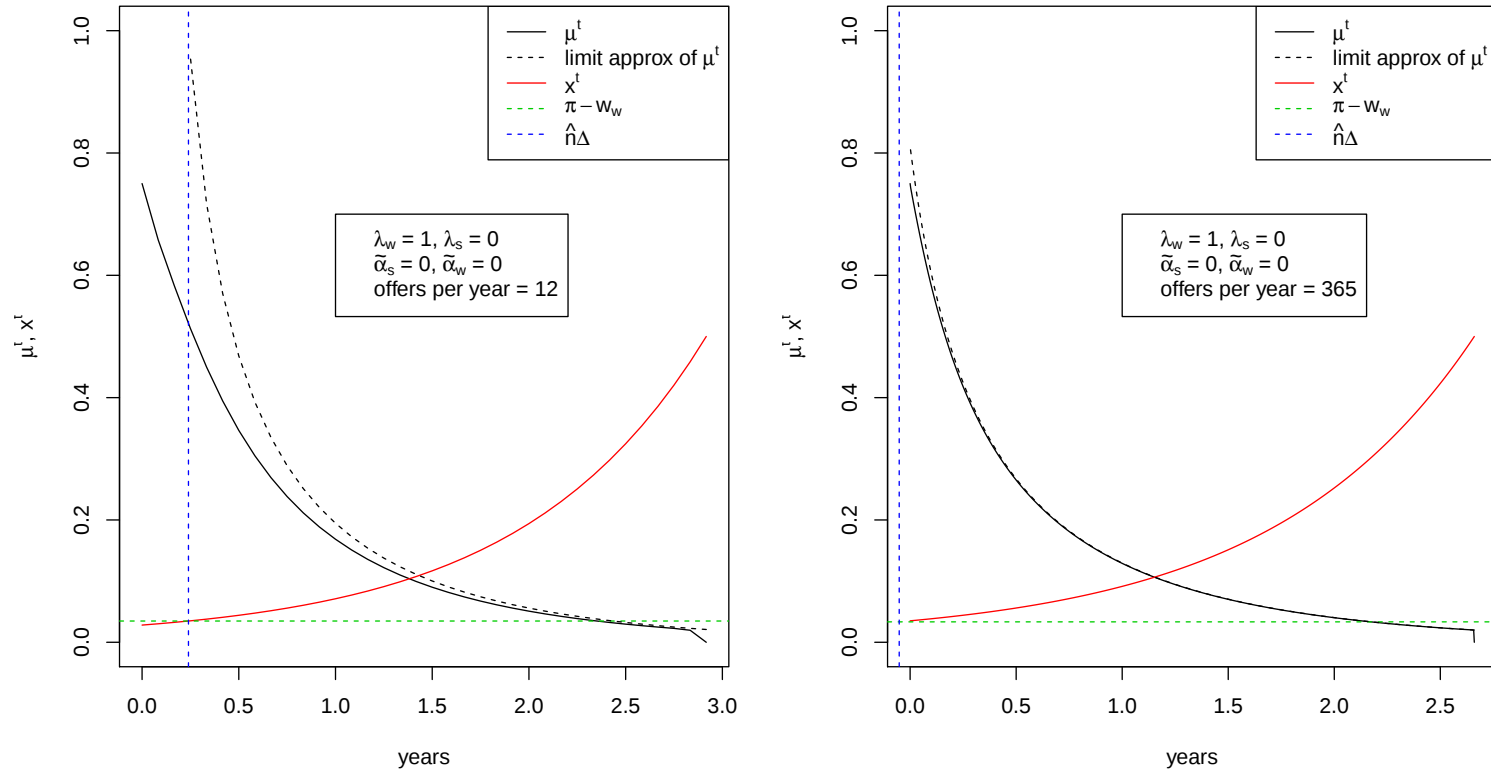
$$\begin{aligned} \frac{o'(\tau)}{o(\tau)} &= \frac{L}{1 - Ae^{\rho'\tau}} \\ \frac{d}{d\tau} \log o(\tau) &= \frac{L}{1 - Ae^{\rho'\tau}} \\ o(\tau) &= \frac{o_1}{1-A} e^{L\tau} \left[\frac{1-A}{1 - Ae^{\rho'\tau}} \right]^{\frac{L}{\rho'}} \end{aligned}$$

The last step is integration, using the initial condition that $o(0) = o^* = o_1/(1-A)$. This initial condition follows directly from lemma 10: For any $\epsilon > 0$ small enough that $Ae^{\rho'\epsilon} < 1$,

$$\frac{o_1}{1-A} < \lim_{\Delta \downarrow 0} o_{\lfloor \epsilon/\Delta \rfloor} < \lim_{\Delta \downarrow 0} O_{\lfloor \epsilon/\Delta \rfloor}(Ae^{\rho'\epsilon}) = \frac{o_1}{1 - Ae^{\rho'\epsilon}} \exp \left(\frac{\epsilon L}{1 - Ae^{\rho'\epsilon}} \right).$$

Figure 3 illustrates convergence of the approximation.

FIGURE 3. Continuous approximation for μ^t with 12 and 365 offers/year



Proposition 3. *The mechanism that maximizes ex ante expected payoffs achieves the first best of certain peace ($n_s = 0$) if either $\pi - \underline{x}_s \geq w_w$ or if the initial belief $o = \mu/(1 - \mu)$ is below a threshold value defined by*

$$o \leq o^M \equiv \frac{\pi - \underline{x}_s - w_s}{w_w - (\pi - \underline{x}_s)}.$$

Otherwise, if $o > o^M$, in the optimal mechanism n_s is the smallest integer such that

$$(17) \quad \frac{\mu}{1 - \mu} \leq \frac{\pi - \underline{x}_s - w_s}{\underline{x}_s - \underline{x}_w} \frac{(\delta\beta_s)^{n_s}}{(\delta\beta_w)^{n_s} - A},$$

and $x_w = (1 - (\delta\beta_w)^{n_s})\underline{x}_w + (\delta\beta_w)^{n_s}\underline{x}_s$ for this n_s .

Proof. Omitted (follows from standard results and algebra using the incentive and participation constraints given in the text).

3. NO-COMMITMENT GAME RESULTS

Proposition 4. *In the no-commitment game with complete information and known rebel type $i = w, s$, for small enough time between offers $\Delta > 0$, there exist subgame perfect equilibria in which the government offers x^* and R_i accepts on the equilibrium path, and the government chooses to implement after R_i accepts. Such equilibria exist for $x^* \in \underline{x}_i \cup [x_i^G, \pi - w_i]$, where $x_i^G = \delta\beta_i\underline{x}_i + (1 - \delta\beta_i)(\pi - w_i)$.*

Proof. We drop the type subscripts $i = s, w$, because they are unnecessary here. So with $x^G \equiv (1 - \delta\beta)(\pi - w) + \delta\beta\underline{x}$, note first that $x^G \in (\underline{x}, \pi - w)$. Also, we clearly cannot support an equilibrium in which $x^* > \pi - w$ is offered and accepted because G would do better by fighting.

Use the following strategies.

- (1) For $x^* = \underline{x}$, G offers $x^t = \underline{x}$ regardless of history and R accepts iff $x^t \geq \underline{x}$ regardless of history. G implements if $x^t \leq x^G$.
- (2) For $x^* \in [x^G, \pi - w]$: Strategies depend on which of two sets of histories of play obtains in period t . In the first, R has never accepted an offer $x^s < x^*$ for any period $s < t$. If so, then in period t , G offers x^* , R accepts any $x^t \geq x^*$, and G implements iff $x^* \leq x^t \leq (1 - \delta\beta)(\pi - w) + \delta\beta x^*$. In the second set of histories, in some period $s \leq t$, R accepted an offer $x^s < x^*$. If $s = t$, then G implements if $x^t < x^G$ and fights otherwise. If $s > t$, then G and R play the strategies in (1).

Case (1) is straightforward. Note also that this equilibrium is Markov Perfect and exists for any $\Delta > 0$.

For case (2), suppose that G deviates to x^t such that $x^G < x^t < x^*$. If R accepts, G 's continuation value for not implementing is $\pi - x^G > \pi - x^t$, so G would fight. Anticipating this, R prefers to reject x^t since $(1 - \delta\beta)\underline{x} + \delta\beta x^* > \underline{x}$. (R would “lose its reputation” by not rejecting the downwards deviation.)

Suppose next that G deviates to $x^t < x^G$. Since G would implement if R accepts, R prefers to reject if $(1 - \delta\beta)\underline{x} + \delta\beta x^* > x^t$, which is certainly true for small enough $\Delta > 0$ since $x^* \geq x^G > x^t$. $\square$

Proposition 5. Consider updating consistent equilibria in the no-commitment game in which $\beta_s > \beta_w > 0$ and the players expect to implement $\underline{x}_i$ in any continuation where G believes that R is certainly type R_i . For small enough time between offers $\Delta > 0$, such equilibria exist and must have the following features:

- On the path, after $n^* - 1$ periods of fighting (if both sides survive this long), G offers a deal x^* where $x^* \in [\underline{x}_s, \pi - w_s]$ and both types of R accept, ending the conflict. There is zero chance of a deal (an accepted, implemented offer) before period $t = n^*$. Offers before then are either rejected by both types or accepted by both types of R , but not implemented by G .
- Maximum war duration, $n^*\Delta$ is increasing in x^* . Thus the shortest feasible maximum conflict duration obtains for the equilibrium with $x^* = \underline{x}_s$. In this equilibrium, n^* is the smallest integer such that

$$n^*\Delta \geq \frac{\log o/o^*}{\lambda_w - \lambda_s},$$

where o^* is the same belief threshold defined in Proposition 2.

Proof. Consider the no-commitment game and restrict attention to updating-consistent equilibria that satisfy the following property:

- P1 The players expect to implement $\underline{x}_i$ in any continuation where G believes that R is certainly type R_i .

Claim 1. In any equilibrium (even without P1) after any history h^t such that $\mu^t \in (0, 1)$, R_s 's continuation payoff for rejecting offer x^t , call it $u^s(\text{No}|x^t, h^t)$ is strictly greater than the weak type's continuation payoff for rejecting, x^t , call it $u^w(\text{No}|x^t, h^t)$.

Proof. R_s can accept any offer than R_w accepts and gets a higher payoff for fighting, $\underline{x}_s > \underline{x}_w$.

Claim 2. R_s accepts an offer $x^t \geq \underline{x}_s$ that G will implement with positive probability in any equilibrium (even without condition P1).

Proof. (i) There is no SGP equilibrium in any continuation with $\mu^t = 0$ in which R_s never accepts.⁴⁶ Thus, if R_s never accepts an offer, it must be that $\mu^t > 0$ for all t on the equilibrium path.

⁴⁶To see this: Suppose to the contrary that there exists a SGP equilibrium in the complete-information game in which R_i never accepts an offer that will be implemented. Then R_i 's payoff is $\underline{x}_i$ and G 's is w_i . But G would definitely implement any accepted offer $x^t = \underline{x}_i + \epsilon \in (\underline{x}_i, x_i^G)$, since it cannot possibly do better by

(ii) It cannot be that both types R_w and R_s never accept any offer x^t on the path that G would implement, because then R_w 's equilibrium payoff is $\underline{x}_w$. G will implement for sure an accepted offer $x \in [\underline{x}_w, x^G]$, since it cannot possibly do better by not-implementing in any equilibrium. R_w would do better to accept $x^t = \underline{x}_w + \epsilon$, and G does better to offer this since for small enough ϵ , $\pi - \underline{x}_w - \epsilon > w_w$.

(iii) Suppose then that R_s never accepts on the path; $\mu^t > 0$ for all t ; and R_w accepts with positive probability some offer that would be implemented on the path. Then R_w must be indifferent between accepting such offers and never accepting (otherwise $\mu^t = 0$ after some amount of time). Thus R_w 's equilibrium payoff would be $\underline{x}_w$. But this cannot occur in equilibrium because, as just argued, G could always and would do better to get acceptance from R_w for sure by offering $x^t = \underline{x}_w + \epsilon$.

(iv) But if R_w accepts for sure in some period t , then $\mu^{t'} = 0$ for $t' > t$, and then it must be that R_s will accept for sure in the continuation, so contradicting the contrary to the claim. $\square$

Claim 3. No separating. (a) If $\mu^t > 0$ and R_s accepts for sure an x^t that would be implemented, then in any P1-equilibrium R_w also accepts x^t . (b) If $\mu^t > 0$, then for small enough $\Delta > 0$, in no P1-equilibrium does R_w accept an offer with probability 1 that R_s rejects.

Proof. (a) Suppose that $\mu^t > 0$ and R_s accepts an x^t that will be implemented, while R_w may reject this x^t . Then following rejection, $\mu^{t+1} = 1$, so by P1 R_w 's payoff is $\underline{x}_w$, which is strictly less than accepting $x^t \geq \underline{x}_s$.

(b) If R_w accepts an offer that R_s rejects, its continuation payoff is at most x_w^G . By deviating to reject (mimicking R_s) it would get $(1 - \delta\beta_w)\underline{x}_w + \delta\beta_w\underline{x}_s$, since $\mu^{t+1} = 0$ implies, with condition P1, that G will offer $\underline{x}_s$. For small enough Δ the latter is strictly greater than x_w^G .

Claim 4. In any P1-equilibrium, there is a period t^ in which, on the equilibrium path, R_s accepts $x^{t^*} \geq \underline{x}_s$ with probability 1. Also, R_s never accepts an offer that would be implemented before this t^* .*

Proof. Claim 2 implies that we only have to show that it cannot be that R_s mixes if offered an x^t that it accepts (and that G would implement) with positive probability on the path.

If $\mu^t = 0$, then by assumption of condition P1 R_s accepts $x^t = \underline{x}_s$ for sure.

If $\mu^t > 0$, then suppose R_s mixes given x^t . R_w must accept for sure, because, using Claim 1,

$$u^w(No|x^t, h^t) < u^s(No|x^t, h^t) = x^t = u^w(Yes|x^t, h^t).$$

not implementing. R_i prefers to accept such an x^t , and G would get, by implementing, $\pi - \underline{x}_i - \epsilon$, which is strictly greater than w_i for small enough $\epsilon > 0$, contradicting the hypothesis.

But this contradicts Claim 3 (no separating). If R_s definitely accepts an offer $x^t \geq \underline{x}_s$ in some period t in any P1-equilibrium, then there is a smallest such t . $\square$

Then we have, immediately,

Claim 5. In any P1-equilibrium, for small enough $\Delta > 0$, on the path R_w never accepts any x^t that would be implemented for $t < t^$, where t^* is the period in which R_s accepts on the path.*

Proof. From Claim 3, there can be no period $t < t^*$ in which, on the path, R_w accepts with probability 1 an offer that would be implemented by G . So we only need to show that it cannot be that R_w could mix on accept and reject in a period $t < t^*$. This would imply that R_w is indifferent. But R_w gets at most x_w^G by accepting, and, for small enough $\Delta > 0$, can get a strictly larger payoff by mimicking R_s and accepting $x^{t^*} \geq \underline{x}_s$ in period $t^* > t$.

The preceding claims sharply delimit the form that a P1 equilibrium can take, making these easy to fully characterize. In any P1-equilibrium, both types of R reject any offer that would be implemented until a period t^* , at which point both accept $x^{t^*} \geq \underline{x}_s$. This implies that G 's beliefs along the equilibrium path are easily characterized by

$$o^t = \frac{\mu^t}{1 - \mu^t} = \left(\frac{\beta_w}{\beta_s} \right)^t \frac{\mu}{1 - \mu}.$$

Using Claims 1-5, suppose that $n^* \geq 0$ is the period in which R accepts $x^{n^*} \geq \underline{x}_s$ for sure, and fix $x^{n^*} = x^* \in [\underline{x}_s, \pi - w_s)$. From the claims, in any P1 equilibrium no offer is accepted and implemented before n^* , so G 's belief evolves according to o^t as just stated. G is willing to implement x^* in period n^* given belief μ^{n^*} if and only if

$$\pi - x^* \geq \mu^{n^*} [(1 - \delta\beta_w)w_w + \delta\beta_w(\pi - x^*)] + (1 - \mu^{n^*})[(1 - \delta\beta_s)w_s + \delta\beta_s(\pi - x^*)],$$

the right-hand side being what G could get by waiting one more period to make the pooling offer x^* .

With algebra, we find that the first period $t = n^*$ at which this constraint is satisfied is the smallest integer t such that

$$o_t = \left(\frac{\beta_w}{\beta_s} \right)^t \frac{\mu}{1 - \mu} \leq \frac{1 - \delta\beta_s}{1 - \delta\beta_w} \frac{\pi - x^* - w_s}{w_w - (\pi - x^*)}$$

$$t\Delta \geq \frac{1}{\lambda_w - \lambda_s} \log \frac{\mu}{1 - \mu} \frac{1 - \delta\beta_w}{1 - \delta\beta_s} \frac{w_w - (\pi - x^*)}{\pi - x^* - w_s}.$$

The right-hand side is increasing in x^* , so that the minimum maximum duration is for $x^* = \underline{x}_s$ as claimed. And it is easy to show that the right-hand side is negative and thus we can support

pooling at time $t = 0$ just when

$$\frac{\mu}{1 - \mu} < \frac{1 - \delta\beta_s}{1 - \delta\beta_w} \frac{\pi - x^* - w_s}{w_w - (\pi - x^*)},$$

which, with $x^* = \underline{x}_s$, is the condition from Proposition 2.

These equilibria can be supported by a large number of off-path beliefs, and a large number of feasible offer sequences by G prior to period n^* . Most straightforwardly, G believes that if R accepts any offer prior to time $t = n^*$, it must be the weak type. G would then implement if the accepted offer was $x^t \in [\underline{x}_w, x^G)$ and not otherwise. $\square$

Proposition 6. *Consider the no-commitment game where $A < 1$ and $\beta_s > \beta_w > 0$. Restrict attention to equilibria that satisfy updating consistency and in which the players expect to implement $\pi - x_w$ in any continuation where G believes that R is certainly the weak type, and G never offers more than $\underline{x}_s$. For small enough time between offers $\Delta > 0$ there exists a $\hat{\mu} > \mu^*$ such that*

- (1) *for initial beliefs $\mu \in (\mu^*, \hat{\mu})$, any such equilibrium is the same as described in Proposition 5;*
- (2) *for $\mu \in (\hat{\mu}, 1)$, any such equilibrium follows the same pattern as in Proposition 1 (the commitment case) for periods $t = 0, 1, \dots, N - \hat{n}$, and then switches to the pattern of Proposition 5 (N is maximum duration less one). $\hat{n}$ is the smallest integer n such that $\pi - w_w \geq ((1 - (\delta\beta_w)^n)\underline{x}_w + (\delta\beta_w)^n\underline{x}_s)$; in the limit $\hat{n}\Delta = \hat{T} = \log(1/A)/(\rho + \lambda_w)$. Thus, G initially makes ascending offers $x^t = x_{N-t}$ which R_s rejects for sure while R_w rejects with a positive probability ν^t until reaching the ex post regret point of $x_{\hat{n}}$. After this G makes any non-serious offers less than $\underline{x}_s$ that R rejects for sure until fighting causes G 's belief to fall to $\mu^t \leq \mu^*$, when G makes the pooling offer $\underline{x}_s$.*

Further $\lim_{\Delta \downarrow 0} \hat{\mu} = \mu^*/(\mu^* + A^{L/\rho'}(1 - \mu^*))$.

Sketch of Proof. Claims 1-4 from the proof of Proposition 5 either hold directly or, for Claims 3 and 4, with minor modification (substituting $\pi - w_w$ for $\underline{x}_w$ in a few places). So in any equilibrium meeting the conditions of the Proposition, there is a period t^* in which R_s accepts $\underline{x}_s$ for sure (and G implements); R_s does not accept any offer that would be implemented before this; and R_w does not accept with probability 1 any prior offer (that would be implemented).

The most R_w can get by accepting an offer that reveals its type before t^* is $\pi - w_w$, since G then can get w_w by going back to war. When there are n periods to go to get to t^* , R_w can get $(1 - (\delta\beta_w)^n)\underline{x}_w + (\delta\beta_w)^n\underline{x}_s$ by mimicking R_s . This is strictly greater than $\pi - w_w$ when $n < \hat{n}$, so R_w surely rejects any offer that would be implemented for $0 < n < \hat{n}$. Thus G 's beliefs evolve in this period only by mechanical learning, specifically,

$$\frac{\mu_n}{1 - \mu_n} = \left(\frac{\beta_s}{\beta_w} \right)^n \frac{\mu_0}{1 - \mu_0},$$

where μ_0 is G 's belief in period t^* , when G offers $x^{t^*} = \underline{x}_s$ and R accepts.

Next, $\mu_0 = \mu^*$, which is the belief such that G is indifferent between ending the game by offering $\pi - \underline{x}_s$, and fighting for one more period and then offer this.⁴⁷ (Obviously μ_0 cannot be greater than this or G would not be willing to pool, and if it is sufficiently smaller then G would have wanted to offer $\underline{x}_s$ earlier than t^* .)⁴⁸

We next describe construction of the equilibrium path for periods $n \geq \hat{n}$. Proofs of uniqueness are essentially the same as for Proposition 1.

Given a sequence of beliefs $\{\mu_n\}_{n=0}^\infty$ with $\mu_0 = \mu^*$ and $\mu_n \geq \psi(\mu_{n-1})$ for $n \geq 1$, define recursively the functions:

$$\begin{aligned} u_n^w &\triangleq (1 - (\delta\beta_w)^n)w_w + (\delta\beta_w)^n(\pi - \underline{x}_s) && \text{for } n < \hat{n} \\ u_n^w(\mu) &\triangleq (1 - \nu_n(\mu))(\pi - x_n) + \nu_n(\mu)((1 - \delta\beta_w)w_w + \delta\beta_w u_{n-1}^w(\mu_{n-1})) && \text{for } n \geq \hat{n} \\ u_n^s &\triangleq (1 - (\delta\beta_s)^n)w_s + (\delta\beta_s)^n(\pi - \underline{x}_s) \\ u_n(\mu) &\triangleq \mu u_n^w(\mu) + (1 - \mu)u_n^s, \end{aligned}$$

where

$$\nu_n(\mu) \triangleq \frac{o(\mu_{n-1})\beta_s}{o(\mu)\beta_w}.$$

The sequence of belief cutoffs are now given by

$$(18) \quad \begin{aligned} \mu_0 &= \mu^* \\ \mu_n &= \psi(\mu_{n-1}) && \text{if } 1 \leq n < \hat{n} \\ u_n(\mu_n) &= u_{n-1}(\mu_n) && \text{if } n \geq \hat{n}. \end{aligned}$$

It is straightforward to verify that this sequence exists and satisfies $\psi(\mu_{n-1}) \leq \mu_n < 1$ for all n ; this is direct for $n < \hat{n}$, and by the same arguments as in Proposition 1 for $n \geq \hat{n}$.

Note that by construction R_w is indifferent between accepting any x_n , $n \geq \hat{n}$, and fighting to $n = 0$ in hopes of getting $\underline{x}_s$. For periods $\hat{n} > n > 0$, R_w strictly prefers fighting to any offer that G would implement if it accepted (off the path). G is likewise indifferent between offering x_n and x_{n-1} when $n > \hat{n}$, by the same construction as in the game with commitment. At $n = \hat{n}$, G is indifferent between offering $x_{\hat{n}}$ and the expected value of the pure fight (till $n = 0$) that will ensue if R rejects.

⁴⁷That is, μ^* from Proposition 2 solves

$$\pi - \underline{x}_s = \mu((1 - \delta\beta_w)w_w + \delta\beta(\pi - \underline{x}_s)) + (1 - \mu)((1 - \delta\beta_s)w_s + \delta\beta_s(\pi - \underline{x}_s)).$$

⁴⁸We ignore here a small range that vanishes in the limit that could allow inconsequential variation in the final belief.

Remarks. (a) For Proposition 6, we have omitted analysis of equilibria in which G pools in the last period on the path on an $x^{t*} > \underline{x}_s$, as discussed in the case analyzed in Proposition 5. Changes are insubstantial but tedious to fully describe. In short, the period of pure fighting lengthens; R_w begins rejecting with probability 1 offers at a period $n' > \hat{n}$.

(b) Propositions 5 and 6 consider only “learning from fighting” environments where $\beta_s > \beta_w$. Table 2 mentions results on limiting maximum and mean durations for the no-commitment game with $\beta = \beta_w = \beta_s$. We briefly describe equilibrium behavior for these somewhat pathological cases here, without giving a full analysis.

Under the assumptions of Proposition 5 (that the players expect to implement $\underline{x}_w$ in any continuation with certainty that $R = R_w$), R_w and G both know that G would implement any accepted offer $x^t \in [\underline{x}_w, x_w^G]$, and is willing to implement $x^t = x_w^G$. For convenience, define ϵ so that $\underline{x}_w + \epsilon = x_w^G$ and note that $\lim_{\Delta \downarrow 0} \epsilon = 0$.

In equilibrium: In the first period, $t = 0$, G offers x_w^G , and R_w rejects with probability $\nu^0 = o^*/o$, so that in $t = 1$ G has belief μ^* and so is indifferent between ending the game by offering $\pi - \underline{x}_s$ and fighting another period before doing so. In $t = 1$ and all subsequent periods, both types of R *always reject* the low offer, and both types accept the high offer. In $t = 1$ and all subsequent periods, G mixes in every period between offering x_w^G and $\underline{x}_s$, putting probability $p \in (0, 1)$ on the high offer, where

$$p = \frac{\epsilon(1 - \delta\beta)}{\delta\beta(\underline{x}_s - \underline{x}_w + \epsilon)}.$$

This p is the probability such that R_w is indifferent between accepting x_w^G and fighting in hopes of getting $\underline{x}_s$ in the next period. (Because R_w is indifferent, R_w is willing to reject with probability 1.)

Because, after the initial offer, R never accepts the low offer, G 's beliefs stay fixed at o^* . There is no limiting maximum duration by which the war ends with probability 1, although the ex ante expected duration is

$$\frac{\mu o^*/o + 1 - \mu}{p + (1 - p)(1 - \beta)} \Delta,$$

which approaches

$$\frac{\mu o^*/o + 1 - \mu}{\lambda} = \frac{1 - \mu}{1 - \mu^*} \frac{1}{\lambda},$$

as Δ gets small.

The Proposition 6 case (players expect to implement $x = \pi - w_w$ if R_w is revealed) is similar, substituting $\pi - w_w$ for $\underline{x}_w$ in the relevant places above. G 's equilibrium high-offer probability p for this case is

$$p = \frac{1 - \delta\beta}{\delta\beta} \frac{\pi - \underline{x}_w - w_w}{w_w - (\pi - \underline{x}_s)}.$$

Computing the limiting mean duration yields

$$\lim_{\Delta \downarrow 0} \frac{\mu o^*/o + 1 - \mu}{p + (1 - p)(1 - \beta)} \Delta = \frac{\mu o^*/o + 1 - \mu}{\lambda + (\rho + \lambda)B} = \frac{1 - \mu}{1 - \mu^*} \frac{1}{\lambda + (\rho + \lambda)B},$$

where $B = \frac{\pi - \underline{x}_w - w_w}{w_w - (\pi - \underline{x}_s)}$.