



Political Brinkmanship and Compromise

Helios Herrera, Antonin Macé, Matias Nùnez

► To cite this version:

Helios Herrera, Antonin Macé, Matias Nùnez. Political Brinkmanship and Compromise. 2025. halshs-03225030v3

HAL Id: halshs-03225030

<https://shs.hal.science/halshs-03225030v3>

Preprint submitted on 11 Feb 2025

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



WORKING PAPER N° 2021-28

Political Brinkmanship and Compromise

**Helios Herrera
Antonin Macé
Matias Nunez**

JEL Codes: C78, D7

Keywords: Brinkmanship, Political Gridlock, Bargaining Advantage.



Political Brinkmanship and Compromise^{*}

Helios Herrera (University of Warwick & CEPR, United Kingdom)

Antonin Macé (CNRS, Paris School of Economics & École Normale Supérieure-PSL,
France)

Matías Núñez (CREST, CNRS, École Polytechnique, GENES, ENSAE Paris, Institut
Polytechnique de Paris, 91120 Palaiseau, France)¹

Abstract

We study how do-or-die threats ending negotiations affect gridlock and welfare when two opposing parties bargain. Failure to agree on a deal in any period implies a continuation of the negotiation. However, under brinkmanship, agreement failure in any period may precipitate a crisis with a small chance. In equilibrium, such brinkmanship threats improve the probability of an agreement, but also increase the risk of crisis. Brinkmanship reduces welfare when one might think it is most needed: severe gridlock. In this case, despite this global welfare loss, a party has incentives to use brinkmanship strategically to obtain a favorable bargaining position.

JEL Classification: C78, D7.

Keywords: Brinkmanship, Political Gridlock, Bargaining Advantage.

^{*}Submitted on October 30, 2023; Revised on September 12, 2024.

¹We greatly benefited from comments and suggestions from Benjamin Blumenthal and Olivier Compte, as well as seminar and conference participants. We are grateful to the editor and three referees for their detailed and constructive feedback. Financial support from the grants ANR-17-CE26-0003, ANR-17-EURE-001 and ANR-11-IDEX-0003/Labex Ecodec/ANR-11-LABX-0047 is gratefully acknowledged.

“Brinkmanship...the threat that leaves something to chance” (Thomas Schelling)

1 Introduction

Motivation. Since the Brexit referendum in June 2016, the exit of the U.K. was marked by everlasting negotiations, with a withdrawal agreement with the E.U. finally ratified by the U.K. in January 2020. For three consecutive times in 2019, the U.K. Parliament rejected a negotiated agreement, which then had to be renegotiated in Brussels. Every time, PM Theresa May, despite threatening not to do so, requested last minute deadline extensions (which were granted by the E.U.) to avoid a no-deal, i.e. a Hard Brexit, that would significantly hurt the U.K. (and E.U.) economies. Indeed, the possibility of a future re-vote on a new deal fostered the unwillingness to compromise of U.K. parties and factions. To try to solve what became known as the “kicking the can down the road problem”, *threats* of no-extension to the ratification deadlines (precipitating a no-deal Hard Brexit outcome) were made on several occasions by E.U. countries.¹ On the U.K. side, in the summer 2019, PM Boris Johnson also vowed for the U.K. to leave the E.U. by Oct 31 2019 “do or die” pledging not to extend the deadline.² Only an early election with a Tory landslide broke the impasse and allowed the withdrawal agreement to be ratified by the U.K.

Similar brinkmanship - “my way or the highway” - episodes characterized political negotiations in other countries, notably in the U.S. in the last decade, where they precipitated several government shutdowns and debt ceiling episodes. For instance, in 2013, the Republican Party in Congress refused to raise the debt ceiling unless President Obama would have defunded the Affordable Care Act (Obamacare), his signature legislative achievement. As negotiations went on, the U.S. Treasury stated that it would have to *delay payments* if funds could not be raised through these measures: the U.S. defaulting on its debt became more likely as days passed without an agreement and would have resulted in permanent damage to the economy. Similar debt ceiling episodes happened in 2011, 2021 and early 2023.³ Although in each case, a U.S. default was avoided, these brinkmanship

¹See for instance, The Guardian Oct 28 2019: “Macron against Brexit extension as Merkel keeps option open”.

²Boris Johnson tried also to suspend the U.K. parliament (provoking a constitutional crisis), thus presenting to the U.K. MPs a Hard Brexit on Oct. 31 as the only alternative to the current agreement. See The Economist Aug. 29 2019: “Taking Back Control”.

³Treasury Secretary Timothy Geithner emphasized the risk posed by these brinkmanship tactics in a 2011 letter

tactics engender a real risk, that is designed to put pressure and extract concessions from the counterparts.

Brinkmanship. Political negotiations or treaty ratifications often occur under soft deadlines (such as reaching a debt ceiling in the US) extendable under certain conditions and/or approval by the negotiating sides. Threatening not to extend these soft deadlines or imposing ‘do-or-die’ conditions for an agreement to be reached, “or else”, may result in a hard breakdown of negotiations, namely an outcome worse than any agreement and, crucially, worse than the status-quo, i.e. the negotiations’ limbo in which no agreement is (yet) reached. It is a brinkmanship tactic which amounts to putting a “time bomb” that can go off anytime over a negotiation. What is key about these brinkmanship threats is that they work by creating the *risk* of an accident. Namely, they generate probabilistic outcomes that hang over the negotiation like a Damocles’s sword that can fall at any moment, albeit with a small chance. Such risks were salient during the U.K. Brexit deal negotiations and U.S. debt ceiling episodes: having substantial effects on financial markets, such brinkmanship episodes represent a *true risk* whose extent depends on many random factors beyond the credibility of the threatening side. As Schelling (1960) observes referring to cold war impasses: “the key to these threats is that, though one may or may not carry them out if the threatened party fails to comply, the final decision is not altogether under the threatener’s control....these risks could involve chance, accident, third-party influence, imperfection in the machinery of decision, or just processes that we do not entirely understand.”

Trade-offs. Political brinkmanship tactics, such as putting a negotiation at the brink of a precipice, raise previously unstudied positive and normative questions which we aim to explore here. Namely, whether imposing such a burden upon the negotiation can be welfare improving, and in particular if a party could benefit from the small chance that a ruinous outcome ends the negotiation. Prima facie, there seems to be two possible benefits of do-or-die threats precipitating calamitous potential costs to all sides. One is a common benefit: making both sides willing to

[Geithner, 2011]: “failure to raise the limit would precipitate a default by the United States. Default would effectively impose a significant and long-lasting tax on all Americans and all American businesses and could lead to the loss of millions of American jobs. Even a very short-term or limited default would have catastrophic economic consequences that would last for decades.”

compromise to find an agreement more urgently thus reducing gridlock, that is reducing the cost of extended negotiations and delayed outcomes. The other is private: do-or-die threats can also be imposed strategically by a negotiating party seeking an advantageous bargaining position, for instance if the blame for a possible Hard Brexit or U.S. default hurts one side more than the other.⁴ On the flip side, there are costs of imposing such threats if the “time bomb” happens to explode before a final agreement is reached and the ruinous outcome de facto materializes, thereby hurting, possibly to a different extent, all parties.

Model. To shed light on the aforementioned trade-offs, we present a model in which two parties must repeatedly decide to approve, or not, a proposed agreement. These proposed agreements are randomly drawn every period out of a set of agreements that grant to both parties (weakly) better outcomes than the status-quo. This randomness reflects the unavoidable underlying uncertainty on the next proposal if the current one fails to be approved.⁵ The core setup we analyze is a standard pie-sharing framework with costly delay, namely the pie is shrinking with time. In the presence of brinkmanship however, if a proposal is rejected, then with some probability $h > 0$ (small, and zero in the benchmark case), this causes the bargaining to end with an outcome *ruinous to both* parties (a crisis). The two sides may differ in their utility from this ruinous outcome, but crucially this utility is always negative, i.e., a *worse outcome than the status-quo*. Conversely, $(1 - h)$ represents the chance that the ruinous outcome is not reached, thus the negotiating process continues through one additional period in which a new proposal is drawn to be voted on, and so forth.

Small Crisis Chance. Our aim is to understand the effects of brinkmanship, namely who may benefit from such a tactic and how this affects several outcomes besides welfare, such as equilibrium gridlock (the per period chance of a deal being rejected), the equilibrium bargaining power (the expected location of an agreement) and the overall equilibrium chance of a ruinous outcome (a Hard Brexit or a U.S. default). A necessary first step to understand the effects of brinkmanship

⁴This share of blame is affected by many political factors beyond the scope of our simple model and ultimately also depends on voters’ perceptions.

⁵Our focus is the final approval of a proposal previously negotiated by a committee/delegation. This negotiation may entail bargaining between several factions as well as unforeseen economic and unanticipated institutional constraints becoming binding. In the case of Brexit, the potential Brexit deals were negotiated externally between the U.K. executive and an E.U. commission and then brought to a vote in the U.K. parliament.

is to disregard the origin or credibility of such threats and to take the value of h as exogenous. We abstract from modeling the intensive margin, i.e. the choice of the value of h (and asking why tactics such as not raising the debt ceiling are credible) as it necessarily implies imposing more structure on the model. A key justification for taking h as exogenous is that in practice, a ruinous outcome such as a U.S. default or a Hard Brexit has a rather small chance of materializing. It is perceived by financial markets as a potentially catastrophic event that has a positive, albeit very low, probability of occurring. In the paper, we rely heavily on the small h assumption to obtain general results. We only focus on the extensive margin in the strategic use of brinkmanship, that is, on the choice of whether to trigger a threat (of known magnitude h) or not by either party.

Results. We find that the unique stationary equilibrium is characterized by an agreement set, representing the scope for agreement, i.e. the deals acceptable by both parties. This agreement set crucially also represents an (inverse) measure of political *gridlock*. We show that, regardless of how ruinous a crisis may be for each side, brinkmanship enlarges the equilibrium scope for agreement, making an agreement more likely in every period. Indeed, it always succeeds in easing gridlock by forcing parties to compromise more overall, although one party may compromise less to extract a bargaining advantage.

The effects of brinkmanship on welfare are subtle. The overarching trade-off is that, while brinkmanship is effective in alleviating gridlock, it is also dangerous: a crisis may realize in equilibrium, which in turn reduces expected welfare. In particular, we show that the effect of brinkmanship threats on welfare crucially depends on the initial level of gridlock (absent brinkmanship), which is determined by parties' patience.

Interestingly, if we start from bargaining positions of severe gridlock (in which disagreement is more likely than not) then enacting a brinkmanship tactic, while easing gridlock, hurts total welfare as the risk of a ruinous outcome outweighs the agreement benefits. Despite ruinous welfare consequences overall in this case, a party can have private benefits from brinkmanship: if one party is advantaged in that it has a lower (perceived political) cost of a crisis, then enacting political brinkmanship shifts the whole agreement set to its advantage. This happens independently of the size of the advantage. If, on the other hand, initial gridlock is mild, then brinkmanship is welfare improving overall and can be promoted by either party if neither is too disadvantaged.

Limitations. This is the first study of how brinkmanship affects several key variables such as gridlock/scope of agreement, the chance of crisis and welfare for both parties at the same time. We want to be upfront about two key assumptions we require to keep the model tractable and obtain these results. First, *exogenous proposals*: this assumption is a stylized modeling device to capture the key idea of *imperfect foresight*; from the point of view of the U.K. House of Commons or U.S. Congress it is impossible to predict, once a deal is voted down, exactly what kind of deal will come to a floor for a future vote, if any. Second, *stationarity*: the model is stationary so as time passes the equilibrium does not change; acceptance thresholds and the chance of agreement/crisis remain constant. We also made this choice for tractability as time-dependent thresholds would be hard to compute, let alone welfare, and would transform the model into a timing game with possibly many equilibria.⁶

In the following, after the literature review, we introduce the general model and its results. To keep the body of the paper compact all proofs are relegated to the appendix.

2 Related Literature

This paper touches on several strands of related literature which we outline below.

War Brinkmanship. Powell [1988] examines “nuclear brinkmanship” in a game of escalation, where the risk of breakdown is endogenously determined and the outcome is binary: either one side or the other wins. Our model is more akin to a bargaining model, as it looks at how brinkmanship changes a continuous bargaining outcome rather than a bang-bang conflict outcome. Schwarz and Sonin [2008] consider brinkmanship in a dynamic setting where a potential aggressor demands concessions from a weaker party under the threat of war and prove that a continuous stream of transfers prevents war more effectively than a lump-sum transfer in the absence of commitment. Cuellar and Rentschler [2023] study commitment strategies in a similar dynamic bargaining environment. In their model, the proposer controls the agenda and decides the intensity of the threat, while we allow for imperfect control of the agenda with given threat intensity. They identify conditions when the proposer commits to starting the conflict to threaten the responder. Acharya and Grillo [2015] study a model of war and peace in which leaders believe there might be crazy types

⁶See the literature review below for a further discussion of these two points.

who always behave aggressively and escalate to a war from a peaceful settlement. This “madman” brinkmanship strategy can be used by non-crazy types to their advantage, as they show. Although we do not model this signalling game, our model yields an outcome similar to that paper: under brinkmanship, a crisis emerges with a positive, albeit small, probability

Political Brinkmanship. Some scholars have studied forms of costly political brinkmanship in the U.S. context. [Patty \[2016\]](#), motivated by U.S. government shutdown episodes, proposes a model in which obstruction in a legislative process is unambiguously costly to all parties due to its consequent delay and gridlock, but incumbents might decide to use this brinkmanship as a signalling device, a show of strength and resolve vis-à-vis voters. [Grillo and Prato \[2020\]](#) propose a theory of democratic backsliding which, as they show, may occur even if citizens intrinsically value democracy. Backsliding, an attack on democratic institutions, can be interpreted as institutional brinkmanship. In a similar spirit to ours, they show how opportunistic authoritarian governments may want to attack democratic institutions to gain an advantage.

Bargaining with Exogenous Offers. Our modeling strategy borrows from the collective search and experimentation models, in which a group chooses every period between accepting the current negotiation outcome or waiting for a (randomly drawn) new outcome next period.⁷ For instance, [Compte and Jehiel \[2010\]](#) show that more stringent majority requirements select more efficient proposals but take more time to do so. [Albrecht et al. \[2010\]](#) find that committees are more permissive than a single decision maker facing an otherwise identical search problem.⁸ [Acharya and Ortner \[2022\]](#) study a dynamic bargaining model with two players and a status quo, where a new alternative is drawn randomly from the set of feasible policies at each step. In their model, players reach a sequence of interim agreements, gradually approaching the Pareto frontier.

Stochastic Bargaining. Several authors assume different types of stochasticity. [Strulovici \[2010\]](#) and [Messner and Polborn \[2012\]](#) focus on committee decisions in which preferences are unknown and only learned over time, thus the option to delay happens in equilibrium albeit with

⁷This literature is somehow related to a classic literature on bargaining where both parties are allowed to search for outside options, see [Wolinsky \[1987\]](#) and [Chikte and Deshmukh \[1987\]](#) for classic treatments on the question. See [Muthoo \[1995\]](#) that analyzes the role of parties being able to leave temporarily the negotiation and [Manzini and Mariotti \[2004\]](#) where bilateral bargaining occurs between parties that can agree on a joint outside option.

⁸In a related model with interdependent values, [Moldovanu and Shi \[2013\]](#) study costly search for a committee and study how acceptance thresholds and welfare depend on the degree of conflict within the committee.

different degrees of efficiency depending on the majority rule. [Ortner \[2017\]](#) analyzes how shocks to the popularity of politicians affect bargaining outcomes. [Moldovanu and Rosar \[2021\]](#) study voting in a Brexit-like model with one irreversible option and compare the effect of different voting rules. They show that voting by supermajority over two consecutive periods dominates voting by simple majority. [Basak and Deb \[2020\]](#) focus on a bargaining environment in which public opinion shocks provide leverage by making compromises costly in the presence of deadlocks.

In our model offers/deals are exogenously stochastically drawn from the set of mutually beneficial outcomes to both parties, this generates naturally inefficient delays in finding agreements. There is some literature on legislative bargaining models with endogenous offers in which elements of stochasticity generate inefficient delays in agreements or gridlock in the presence of an endogenous status-quo. Several papers analyze stochastically evolving preferences, see [Dziuda and Loeper \[2016\]](#) or [Bowen et al. \[2017\]](#). Other works explore the case of delay with a stochastic total surplus, such as [Merlo and Wilson \[1995\]](#), [Merlo and Wilson \[1998\]](#) and [Eraslan and Merlo \[2002\]](#).

Bargaining with Deadlines. Several authors have looked at the effect of hard deadlines in negotiations, which nicely complements our stationary setup. In other words, while we study dynamic negotiation between two parties in the presence of a stationary stochastically extendable deadline, in most of the literature, the deadline is tight in the sense that no extension is possible.⁹ This generates incentives to reach agreements in the “eleventh hour”, that is at or very close to the deadline (see [Simsek and Yildiz \[2016\]](#) for the role of optimism in these models). Such (non-stationary) games have been studied by [Ma and Manove \[1993\]](#), [Fuchs and Skrzypacz \[2013\]](#) and others.¹⁰

⁹See also [Ellingsen and Miettinen \[2008\]](#) who study a model of bilateral bargaining where negotiators can write binding contracts and show that conflict is frequently the unique equilibrium outcome when commitments technologies are highly credible.

¹⁰See [Cramton and Tracy \[1992\]](#) for empirical evidence or [Gneezy et al. \[2003\]](#) for experimental evidence on this observation.

3 Model

3.1 Setup

There are two parties, denoted by $\theta = 0$ and $\theta = 1$. Each party $\theta \in \{0, 1\}$ has single-peaked preferences over $X = [-1, 1]$, the set of all possible *deals*. Specifically, we assume that $u_0(x) = 1 - x$ and $u_1(x) = 1 + x$, so that party 0 (resp. 1) prefers lower (resp. higher) deals.

The final outcome may be a deal in X or the *ruinous outcome (or crisis)* d^* . The option d^* does not lie in X and yields each party $\theta \in \{0, 1\}$ a utility $d_\theta < 0$. We denote by $D = d_0 + d_1 < 0$ the ruinous outcome's overall *severity*, and by $B = d_1 - d_0$ party 1's *advantage* in terms of ruinous outcome payoffs relative to party 0. This advantage could be actual or simply perceived: e.g. represent a perceived political advantage based on which side voters would blame more if a crisis materializes. Without loss of generality, we assume that $B \geq 0$, and if $B > 0$, we refer to 1 as the *advantaged* party and to 0 as the *disadvantaged* one. The main parameters of the model are represented on [Figure 1](#).

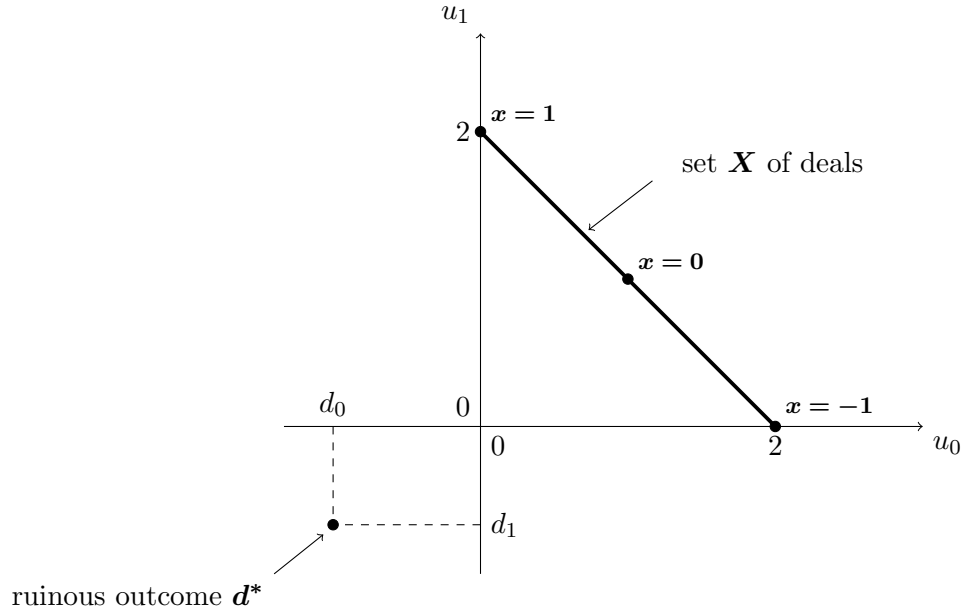


Figure 1: Utilities in the Bargaining Model

In order to focus on the effects of the ruinous outcome payoffs on delay/gridlock and welfare, we assume a simple bargaining protocol which naturally allows for some stochastic delay in equilibrium. The bargaining procedure is sequential and the agenda has a stochastic aspect: for each $t \in \mathbb{N}$,

a proposed deal x_t is drawn, from the set of deals X , uniformly and independently of previous draws.¹¹ Then, parties simultaneously choose to accept or reject the proposal x_t . If both accept it at period t , the final outcome is x_t . Otherwise, the resulting outcome is governed by a binary outcome H_t which takes value 1 with a small probability h , and value 0 with probability $1 - h$. If $H_t = 1$, a crisis occurs: the ruinous outcome d^* is obtained at period t and the game ends. If, on the contrary, $H_t = 0$, then in that period the crisis is averted and negotiations continue with both parties moving to the next period $t + 1$. The parameter $h \in [0, 1)$ thus represents the *brinkmanship threat*, the chance the ruinous outcome materializes in any period in which disagreement persists.¹² In the benchmark case there is no brinkmanship, thus $h = 0$. Even when brinkmanship is present, we will assume that the threat h is small. This is the interesting and empirically relevant case in both U.S. debt ceiling episodes and Brexit negotiations.¹³ Both parties share the same discount factor $\beta \in (0, 1)$.

3.2 Stationary Equilibrium

The strategy of each party consists in accepting or rejecting deals as they arrive. We restrict our attention to stationary strategies. For a given (stationary) strategy profile, we denote by $A \subseteq X$ its *agreement set*, i.e. the set of deals that would be accepted by both sides if proposed. We denote by w_θ party θ 's *reservation value*, i.e. her expected continuation utility when she rejects a deal. By stationarity, w_θ satisfies the following recursive equation:

$$w_\theta = \overbrace{hd_\theta}^{\text{ruinous outcome } d^*} + \overbrace{\beta(1-h)\mathbb{P}(x \in A)\mathbb{E}[u_\theta(x) \mid x \in A]}^{x \text{ lies in the agreement set}} + \overbrace{\beta(1-h)\mathbb{P}(x \notin A)w_\theta}^{x \text{ does not lie in the agreement set}}.$$

The previous formula can be read as follows : when a deal is rejected, the ruinous outcome arises with probability h ; with complementary probability $(1 - h)$ the game continues to the next stage,

¹¹The reduced-form assumption that proposals are randomly generated provides a tractable way to account for the inevitable uncertainty surrounding future proposals. This assumption is common in the bargaining literature, see for instance Penn [2009], Compte and Jehiel [2010] and Acharya and Ortner [2022]. The former article provides a detailed discussion of this assumption.

¹²Such an exogenous risk of breakdown also appears in one of Binmore et al. [1986]'s approaches to relate strategic bargaining models to the Nash bargaining solution.

¹³We discuss in Appendix A.2 the results obtained for intermediate and large threats.

where payoffs are discounted by β ; in that stage, either the proposal x is accepted and the expected payoff is that of an average deal in the agreement set A , or the proposal x is rejected, in which case the payoff is the reservation value w_θ .

We assume that each party accepts a deal if and only if her payoff is greater than or equal to her reservation value w_θ , even if her opponent rejects it (thus making her a priori indifferent between accepting and rejecting).¹⁴ Then, the condition for a stationary profile with reservation values (w_0, w_1) to be an equilibrium is that its agreement set satisfies $A = A_w = \{x \in X \mid u_\theta(x) \geq w_\theta, \forall \theta \in \{0, 1\}\}$.¹⁵ Hence, this profile is an equilibrium if and only if

$$w_\theta = \frac{hd_\theta + \beta(1-h)\mathbb{P}(x \in A_w)\mathbb{E}[u_\theta(x) \mid x \in A_w]}{1 - \beta(1-h)\mathbb{P}(x \notin A_w)}. \quad (1)$$

The key (endogenous) outcome variables of our analysis are the following ones. We denote by $\lambda = \frac{|A|}{|X|} = \frac{|A|}{2}$, where $|\cdot|$ denotes the Lebesgue measure, the *agreement probability* in every period. Importantly, the variable $(1 - \lambda)$ also represents the level of equilibrium *gridlock* in the legislature. We denote by c the *center* of the agreement set A , which is a measure of the relative bargaining strength of each side in equilibrium: the farther c is from zero the stronger is the equilibrium bargaining power of one party relative to the other. With these notations, the agreement set can be written as $A = [c - \lambda, c + \lambda]$, as illustrated on [Figure 2](#) below.

¹⁴This is equivalent to ruling out weakly dominated strategies (in the stage game), a standard assumption in bargaining. This assumption can also be thought as being similar to the one of partial honesty used in mechanism design (see [\[Dutta and Sen, 2012\]](#) among others). The role of this assumption is simply to rule out equilibria where both parties simultaneously reject a deal that they both strictly prefer to the outcome they obtain in case of rejection.

¹⁵Note that the monotonicity of the utility functions implies that A_w is an interval.

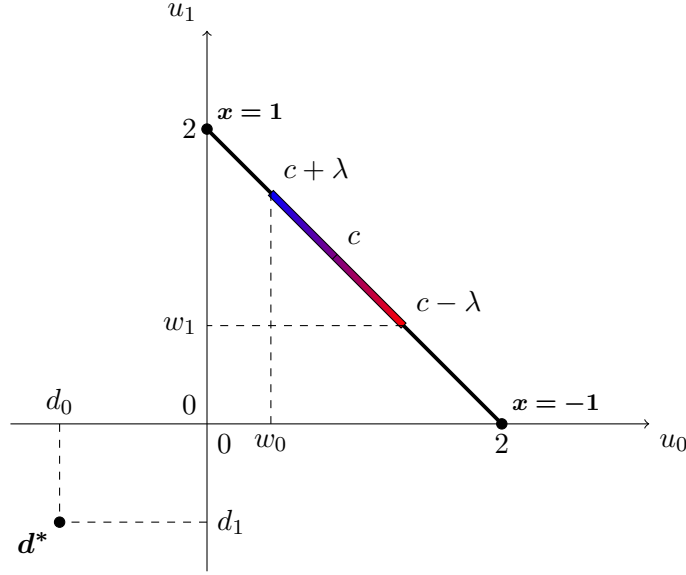


Figure 2: Agreement set in the Bargaining Model

4 Gridlock and Welfare

4.1 Agreement set

In this section, we fully characterize the (unique) equilibrium agreement set when the brinkmanship threat h is small enough, as specified below. In this case the equilibrium is *interior*, namely its agreement set is strictly contained in the proposal set X , or equivalently both agents have positive reservation values, i.e. $w_0, w_1 > 0$.

Theorem 1 *There exists $h_0 > 0$ such that for any $h < h_0$, there exists a unique stationary equilibrium and this equilibrium is interior.*

Henceforth, we restrict our analysis to $h < h_0$, i.e. to the cases where h is either zero (benchmark case of no brinkmanship) or small enough when positive.

Proposition 1 *The equilibrium agreement probability λ is increasing in the brinkmanship threat h . If party 1's advantage B is positive, then, in equilibrium, the relative bargaining power c is increasing in the brinkmanship threat h .*

To illustrate [Proposition 1](#), we display on [Figure 3](#) the upper and lower bounds and relative bargaining power i.e. the center c of the agreement set as functions of the threat h , in a numerical

example. The blue (resp. red) curve represents the worst deal that the disadvantaged (resp. advantaged) party is willing to accept, which coincides with the upper bound (resp. lower bound) of the agreement set.

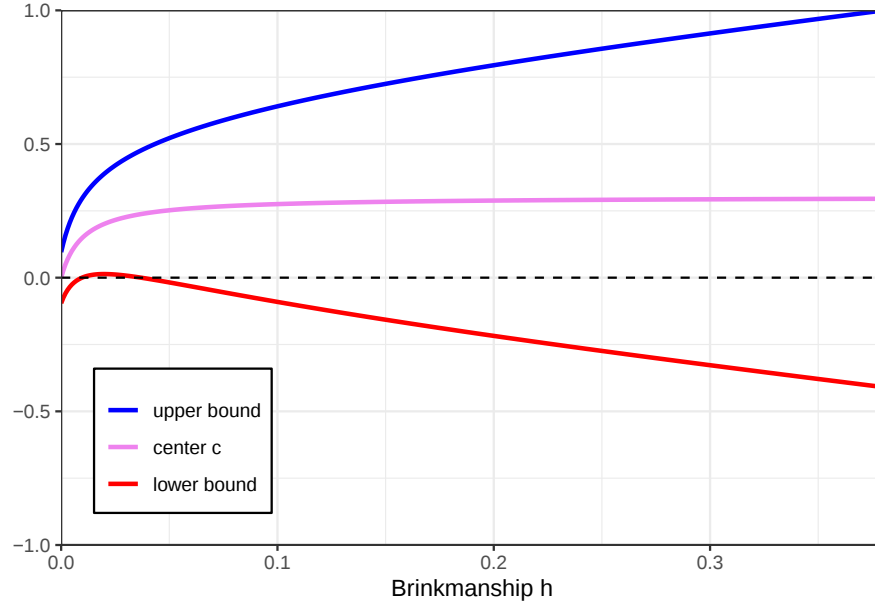


Figure 3: Agreement set for $\beta = 0.99$, $D = -1$ and $B = 0.6$.

The first key point from [Proposition 1](#) is that *brinkmanship eases gridlock*: the scope for agreement, or overall willingness of both parties to accommodate the other party, as measured by the agreement probability λ , always increases with the brinkmanship threat h . Note that, although parties compromise more overall under brinkmanship, it could be that one party compromises less. This is indeed what happens on [Figure 3](#) for low values of h , as party 1 compromises less to push her bargaining advantage.

Second, the relative bargaining power in equilibrium represented by the center of the agreement set c , namely the deal outcome obtained on average, exhibits a monotonic pattern. Starting from the benchmark of no brinkmanship, c increases when brinkmanship is turned on, making the advantaged party better-off as long as a deal is agreed upon before a crisis ensues.

We now describe how, for a given level of threat h , the characteristics of the ruinous outcome (its severity D and asymmetry B) affect the bargaining power and gridlock.

Proposition 2 *At equilibrium, the relative bargaining power c and agreement probability λ are*

such that

$$c = \frac{\Phi B}{2} \quad \text{and} \quad \lambda = \frac{1}{2\Delta} \left(\sqrt{1 + 4\Delta \left(1 - \frac{\Phi D}{2}\right)} - 1 \right), \quad (2)$$

where $\Phi > 0$ and $\Delta > 0$ only depend on brinkmanship threat h and discount factor β .

Firstly, for a given threat h , the equilibrium bargaining advantage c only depends on the payoff advantage in case of ruinous outcome B . In particular, conditional on an agreement, a deal is closer to party 1's bliss point when this advantage is larger. Second, the level of gridlock or the agreement probability λ only depends on the severity D : when the ruinous outcome is more severe ($|D|$ is higher), parties are more likely to compromise overall and gridlock is less severe.

4.2 Joint Welfare

We now characterize parties' equilibrium joint welfare, to try to understand if and when brinkmanship is ever beneficial overall. Denoting by W_θ party θ 's expected welfare, our key outcome variable here is $W = \frac{W_0 + W_1}{2}$ the *joint welfare*, which proxies the total average benefit from the expected bargaining outcome. Welfare W_θ satisfies the following recursive equation:

$$W_\theta = \overbrace{\mathbb{P}(x \in A) \mathbb{E}[u_\theta(x) \mid x \in A]}^{x \text{ lies in the agreement set}} + \overbrace{\mathbb{P}(x \notin A) (hd_\theta + \beta(1-h)W_\theta)}^{x \text{ does not lie in the agreement set}}.$$

The welfare W_θ is determined at the beginning of a period. Thus, either the randomly selected deal x belongs to the agreement set, in which case it yields an expected utility $\mathbb{E}[u_\theta(x) \mid x \in A]$, or it fails to do so and hence, either the ruinous outcome is selected or a new period starts, with expected welfare W_θ . Therefore, party θ 's welfare is given by:

$$W_\theta = \frac{(1-\lambda)hd_\theta + \lambda \mathbb{E}[u_\theta(x) \mid x \in A]}{1 - \beta(1-h)(1-\lambda)}. \quad (3)$$

The following result describes how brinkmanship affects parties' joint welfare. A key parameter in this result is the initial agreement probability $\lambda^0 := \lambda(h=0)$, which is determined by the discount factor β .

Theorem 2 *The effect of brinkmanship on welfare depends on parties' patience:*

- When $\beta > 2/3$ (patient parties), the initial gridlock is severe ($\lambda^0 < 1/2$) and brinkmanship hurts parties' joint welfare.
- When $\beta < 2/3$ (impatient parties), the initial gridlock is mild ($\lambda^0 > 1/2$) and brinkmanship improves parties' joint welfare.

In both cases, the effect of brinkmanship on welfare is more pronounced when the crisis is more severe ($|D|$ is higher), and it is independent of party 1's advantage (B).

Theorem 2 states that the effect of brinkmanship on welfare crucially depends on parties' patience, or equivalently, on the initial severity of gridlock (absent brinkmanship).¹⁶ When gridlock is mild to begin with ($\lambda^0 > 1/2$) then brinkmanship is desirable. But, more interestingly, the converse can also be true: political brinkmanship hurts precisely when it is most needed to break the gridlock. Put differently, when disagreement is large to begin with ($\lambda^0 < 1/2$), brinkmanship is bad for welfare overall because the negative effect of a higher likelihood of crisis outweighs the gain from the enlarged agreement scope. In sum, it is too risky to use brinkmanship when initial gridlock is severe, and the result further asserts that this negative effect is magnified by the severity of a crisis.

4.3 Intuition for Joint Welfare Result

Here we shed light on the forces which govern the link between brinkmanship and welfare. We decompose parties' joint welfare as the product of two factors: *instant welfare* and *timing factor*.

Proposition 3 *Parties' joint welfare can be expressed as:*

$$W = \underbrace{\left(1 - \mathbb{P}(d^* | A) + \left(\frac{D}{2}\right) \mathbb{P}(d^* | A)\right)}_{\text{instant welfare}} \times \underbrace{f(h)}_{\text{timing factor}},$$

where $\mathbb{P}(d^* | A) := \frac{(1 - \lambda)h}{1 - (1 - \lambda)(1 - h)}$ denotes the overall risk of crisis and the timing factor $f(h)$ is increasing in h .

¹⁶As $\lambda^0 = \frac{2}{1 + \sqrt{1 + 4\beta/(1 - \beta)}}$ (see (13) derived in the proof of **Theorem 2**) is a decreasing function of β , the condition that λ^0 is small enough is equivalent to the requirement that the discount factor β is high enough.

The instant welfare is the undiscounted welfare of the *eventual* bargaining outcome. As the average utility of an accepted bargaining deal is 1, the instant welfare is a convex combination of 1 and $\frac{D}{2}$, weighted by the overall risk of crisis $\mathbb{P}(d^* | A)$. The timing factor accounts for the *expected delay* incurred in reaching an agreement through the bargaining process.

The shape of welfare as a function of the threat h thus depends on the relative importance of those two terms. The timing factor is increasing in the threat h , as parties reach an agreement sooner with brinkmanship. By contrast, the instant welfare decreases with h , as the overall risk of crisis $\mathbb{P}(d^* | A)$ is null for $h = 0$ but becomes positive for $h > 0$. To provide intuition on [Theorem 2](#), we may write:

$$\left. \frac{\partial W}{\partial h} \right|_{h=0} = f'(0) + f(0) \left(-1 + \frac{D}{2} \right) \left. \frac{\partial \mathbb{P}(d^* | A)}{\partial h} \right|_{h=0} = f'(0) + f(0) \left(-1 + \frac{D}{2} \right) \frac{1 - \lambda^0}{\lambda^0}.$$

We see here that the positive effect of brinkmanship on welfare (through the timing factor) is constant, while its negative effect (through the risk of crisis, or equivalently, the instant welfare) decreases with the agreement probability λ^0 (i.e. it increases with the initial gridlock level). Hence, brinkmanship becomes welfare decreasing when the initial gridlock is too severe.

4.4 Welfare divergence

In this short section, we discuss how brinkmanship affects the welfare divergence between the two parties $W_1 - W_0$.

Proposition 4 *The welfare divergence is given by $W_1 - W_0 = 2c = \Phi B$, where the scale factor $\Phi := \frac{h}{1 - \beta(1 - h)}$ is increasing with brinkmanship h .*

The result of [Proposition 4](#) is intuitive: the welfare divergence increases with brinkmanship h , while being proportional to the crisis bias B . The scale factor Φ is the fixed point of the equation $Y = h + (1 - h)\beta Y$. Thus, Φ corresponds to the (discounted) expectation of a lottery yielding 1 whenever a crisis occurs, in a benchmark where parties would reject all proposals. In turn, Φ can be interpreted as a refined discount factor, which incorporates both the benchmark discount factor β and the brinkmanship threat h .

5 Strategic Use of Brinkmanship

In the previous section we analyzed the effects of political brinkmanship and explored when brinkmanship can be jointly welfare improving. In this section, we take a step back and focus on whether political brinkmanship may benefit one particular party (compared to the benchmark of no brinkmanship), and may thus be used strategically to gain a bargaining advantage by one or either of the two parties. As before, we consider the empirically relevant case in which the brinkmanship threat h is small and positive. We thus say that party θ *has an incentive to use brinkmanship* if $\left. \frac{\partial W_\theta}{\partial h} \right|_{h=0} > 0$.

For instance, as in our salient example in the introduction, in the US the two main parties can have budget negotiations with or without lifting the debt ceiling first: if one side disagrees to lift the debt ceiling then evidently this party is forcing the negotiation to occur under brinkmanship.

5.1 Brinkmanship by Advantaged Party

We focus first on the advantaged side, party 1. When, absent brinkmanship, gridlock is mild, i.e. $\lambda^0 > 1/2$, we know from [Theorem 2](#) that brinkmanship is welfare-improving overall. Clearly in that case, since brinkmanship further creates a welfare gap advantaging party 1 ([Proposition 4](#)), this party has an incentive to use brinkmanship strategically. The following result provides sufficient conditions under which party 1 will be willing to use brinkmanship strategically (e.g. not raise the U.S. debt ceiling), even when it is welfare-decreasing overall.

Theorem 3 *Party 1 has an incentive to use brinkmanship in each of the following cases:*

1. *for any advantage $B > 0$, if parties are patient enough, that is, when initial gridlock is severe enough (i.e. sufficiently high β , or equivalently, small λ^0),*
2. *for any (sufficiently severe) ruinous outcome $D < -1/8$, if the advantage B is high enough.*

[Theorem 3](#) emphasizes that there are several grounds for the strategic use of brinkmanship by party 1 even though it decreases welfare overall. One sufficient condition is that initial gridlock be high enough (which happens when parties are sufficiently patient), and it holds for any severity and any asymmetry of the ruinous outcome. Indeed, when initial gridlock is high, the private bargaining advantage obtained by party 1 (location of an expected deal, i.e. c) is an order of

magnitude higher than the common drop in welfare. The second condition states that whenever the ruinous outcome's severity is not negligible ($D < -1/8$), enough asymmetry between the two sides in ruinous outcome payoffs ensures that the advantaged party profits from engaging in brinkmanship tactics. The reason is that the bargaining concessions made by the disadvantaged party 0 in equilibrium are sufficient to overturn the drop in welfare.

The welfare consequences of the brinkmanship strategy are illustrated in [Figure 4](#).

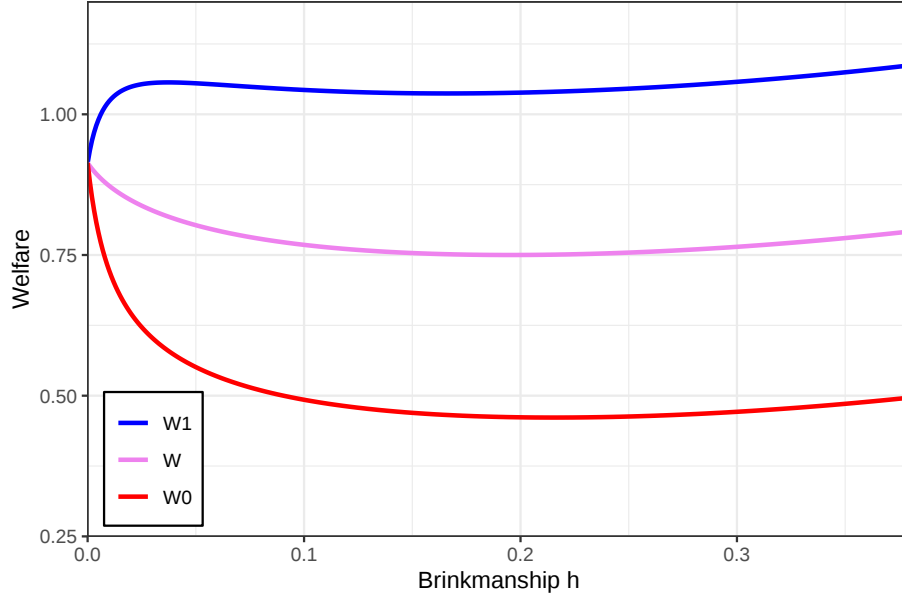


Figure 4: Parties' welfare for $\beta = 0.99$, $D = -1$ and $B = 0.6$

In the example depicted in [Figure 4](#), absent brinkmanship gridlock is severe: the agreement probability in any given period is $\lambda^0 \approx 10\%$. Hence, a small brinkmanship threat is welfare-decreasing overall, but the brinkmanship strategy is nevertheless advantageous for party 1.

5.2 Brinkmanship by Disadvantaged Party

The brinkmanship strategy is never profitable for Party 0 in situations of intense initial gridlock, i.e. $\lambda^0 < 1/2$, as illustrated on [Figure 4](#). In that case, brinkmanship deteriorates overall welfare ([Theorem 2](#)) and also creates a welfare gap to the disadvantage of party 0 ([Proposition 4](#)). Yet, in some situations which we outline below even the disadvantaged party 0 may want to use brinkmanship.

Proposition 5 *Party 0 has an incentive to use brinkmanship when parties are sufficiently impatient ($\beta < 2/3$), that is, when initial gridlock is mild ($\lambda^0 > 1/2$), and its disadvantage B is small enough.*

The intuition for [Proposition 5](#) is that absent significant initial gridlock, brinkmanship improves the willingness to compromise of parties so much that even the disadvantaged party may benefit from it. We illustrate on [Figure 5](#) below that the positive welfare consequences for the disadvantaged party only accrue when the disadvantage is small enough.¹⁷ Note that in this case, the brinkmanship strategy is Pareto-improving.

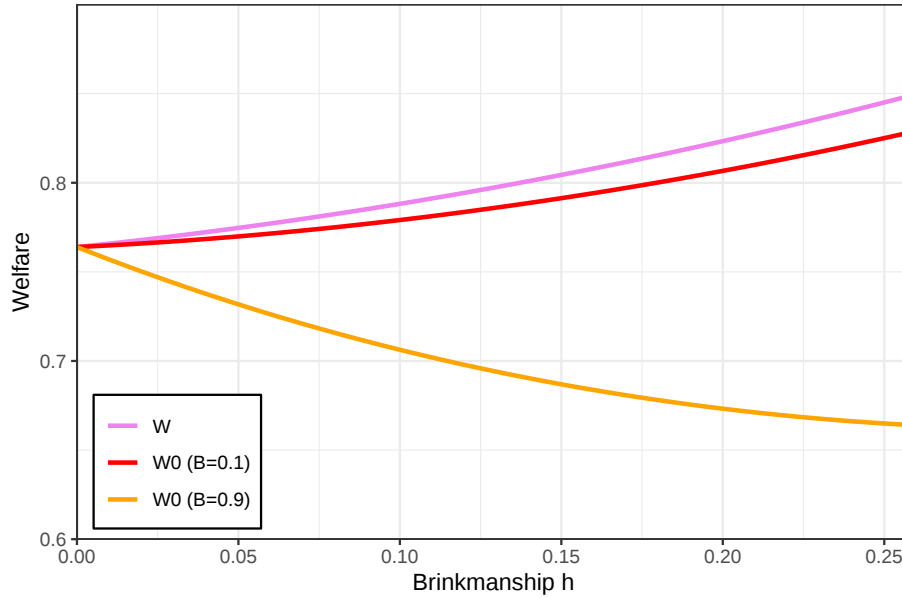


Figure 5: Joint and party 0's welfare for $\beta = 0.5$ and $D = -1$

6 Conclusion

The broad question of how political brinkmanship tactics during political negotiations may alter agreement outcomes seems salient in the current polarized political climate, though it is theoretically still largely unexplored. Here we analyzed this in a simple model that features some disagreement by default, namely stochastic agreement delay. In particular, we studied how brinkmanship affects

¹⁷Note that the joint welfare W does not depend on B , by application of [Lemma 1](#) in [Appendix A.1](#) (W only depends on λ) and [Proposition 2](#) (λ does not depend on B).

gridlock/delay, the welfare of the negotiating parties and finally when parties have incentives to use such brinkmanship tactics for expected political gains.

The key takeaway is that in situations of severe gridlock political brinkmanship hurts global welfare of the negotiating parties as it raises too much the risk of a crisis. However, even in these situations of intense gridlock one party has incentives to use this brinkmanship tactic to gain a bargaining advantage in the negotiation. In sum, our theoretical framework allows us to disentangle shortcomings but also advantages of brinkmanship, a recurring political tactic, and illustrates precisely when one party can use it profitably, hurting the bottom line overall.

The model could be extended in several directions. For tractability reasons, we maintained the assumptions of linear utilities, homogeneous discount factor and uniform distribution of deals. It would be worthwhile to explore relaxations of these assumptions. One could also consider introducing asymmetric information, in which the payoffs to ruinous outcomes would be private information of the two sides. Finally, we avoided credibility issues, assuming small brinkmanship threats, which seem empirically relevant. But the results of our model may be used as the last stage of a broader model which endogenizes the credibility of walk-away do-or-die threat announcements.

A Appendix

A.1 Proofs

We introduce two parameters that will be used throughout the proofs. We note $\Phi(h) = \frac{h}{1-\beta(1-h)}$ and $\Delta(h) = \frac{\beta(1-h)}{1-\beta(1-h)}$. Observe that Φ and Δ are continuous functions of h , that Φ increases with h , while Δ decreases with h and that $\Phi(h=0) = 0$ and $\Delta(h=0) = \frac{\beta}{1-\beta}$. When there is no confusion, we simply write Φ for $\Phi(h)$, and Δ for $\Delta(h)$.

A.1.1 Proof of **Theorem 1**, **Proposition 1** and **Proposition 2**

Proof.

We classify stationary strategy profiles into three types: *no-compromise (or interior)* if $w_0, w_1 > 0$; *partial-compromise* if $(w_0 \leq 0 \text{ and } w_1 > 0)$ or $(w_0 > 0 \text{ and } w_1 \leq 0)$; and *full-compromise* if $w_0, w_1 \leq 0$. In the sequel, we show that for h small enough, only interior equilibria exist.

1. Reservation values. We start by computing parties' reservation values as functions of the

agreement set (at a given stationary profile). We apply (1), by letting $\lambda := \mathbb{P}(x \in A)$ and observing that, since $u_1(x) = 1 + x$, we have $\mathbb{E}[u_1(x) \mid x \in A] = u_1(c) = 1 + c$. We obtain the reservation value for party 1:

$$w_1 = \frac{hd_1 + \beta(1-h)\lambda(1+c)}{1 - \beta(1-h) + \beta(1-h)\lambda} = \frac{\Phi d_1 + \Delta\lambda(1+c)}{1 + \Delta\lambda}. \quad (4)$$

Similarly, as $u_0(x) = 1 - x$, we have $\mathbb{E}[u_0(x) \mid x \in A] = u_0(c) = 1 - c$. By application of (1), we obtain:

$$w_0 = \frac{\Phi d_0 + \Delta\lambda(1-c)}{1 + \Delta\lambda}. \quad (5)$$

2. Agreement sets in equilibrium.

No-compromise equilibrium. Let w be a no-compromise equilibrium, i.e. such that $w_0 > 0$ and $w_1 > 0$. The agreement set is thus $A = [c - \lambda, c + \lambda] = [-1 + w_1, 1 - w_0]$. We obtain :

$$\begin{cases} w_1 + w_0 &= 2(1 - \lambda) \\ w_1 - w_0 &= 2c. \end{cases}$$

Solving for c first, we get:

$$2c = \frac{\Phi(d_1 - d_0) + \Delta\lambda 2c}{1 + \Delta\lambda} \Leftrightarrow c = \frac{\Phi(d_1 - d_0)}{2} = \frac{\Phi B}{2}. \quad (6)$$

Solving for λ , we obtain :

$$2(1 - \lambda) = \frac{\Phi(d_0 + d_1) + 2\Delta\lambda}{1 + \Delta\lambda} = \frac{\Phi D + 2\Delta\lambda}{1 + \Delta\lambda} \Leftrightarrow 1 - \frac{\Phi D}{2} = \lambda + \Delta\lambda^2. \quad (7)$$

Hence, the agreement probability equals:

$$\lambda = \frac{1}{2\Delta} \left(\sqrt{1 + 4\Delta \left(1 - \frac{\Phi D}{2} \right)} - 1 \right). \quad (8)$$

Partial-compromise equilibrium ($w_0 \leq 0$ and $w_1 > 0$). Let w be a partial-compromise equilibrium with $w_0 \leq 0$ and $w_1 > 0$. Then we have $A = [-1 + w_1, 1]$, so that some proposals, that are too far from party 1's bliss point, are rejected. We have $w_1 = 2(1 - \lambda)$ and $1 + c = 1 + \frac{w_1}{2} =$

$2(1 - \frac{\lambda}{2})$. Using (4), we obtain:

$$2(1 - \lambda)(1 + \Delta\lambda) = \Phi d_1 + 2\Delta\lambda \left(1 - \frac{\lambda}{2}\right) \Leftrightarrow 1 - \frac{\Phi d_1}{2} = \lambda + \Delta \frac{\lambda^2}{2}.$$

Hence, the agreement probability equals:

$$\lambda = \frac{1}{\Delta} \left(\sqrt{1 + 2\Delta \left(1 - \frac{\Phi d_1}{2}\right)} - 1 \right).$$

The center of the agreement set is given by $c = 1 - \lambda$.

Full-compromise equilibrium. In a full-compromise equilibrium ($w_0 \leq 0$ and $w_1 \leq 0$), the agreement set is $A = X = [-1, 1]$.

3. Conditions for existence. We first provide conditions for the existence of each type of equilibrium. In particular, we observe that partial-compromise equilibria with $w_0 > 0$ and $w_1 \leq 0$ cannot exist.

No-compromise equilibrium: existence. A no-compromise equilibrium is characterized by the system of equations: $c = \frac{\Phi B}{2}$ and $1 - \frac{\Phi D}{2} = \lambda + \Delta \lambda^2$. As $c = \frac{\Phi B}{2} \geq 0$, a necessary and sufficient condition for such an equilibrium to exist is that the previous system admits a solution with $c + \lambda < 1$, or equivalently $\lambda < 1 - \frac{\Phi B}{2}$. Hence, a no-compromise equilibrium exists if and only if:

$$\exists \lambda \in \left[0, 1 - \frac{\Phi B}{2}\right), \quad 1 - \frac{\Phi D}{2} = \lambda + \Delta \lambda^2. \quad (9)$$

Note that $\lambda \mapsto \lambda + \Delta \lambda^2$ is increasing and continuous on $[0, 1 - \frac{\Phi B}{2})$, and takes value $0 < 1 - \frac{\Phi D}{2}$ for $\lambda = 0$. Hence, by the intermediate value theorem, condition (9) is equivalent to $1 - \frac{\Phi D}{2} < 1 - \frac{\Phi B}{2} + \Delta(1 - \frac{\Phi B}{2})^2$. This condition can then be written: $\frac{\Phi(B-D)}{2} < \Delta(1 - \frac{\Phi B}{2})^2$. Let $F_0(h) = \frac{\Phi(h)(B-D)}{2} - \Delta(h) \left(1 - \frac{\Phi(h)B}{2}\right)^2$. To sum up, we have obtained that a no-compromise equilibrium exists if and only if $F_0(h) < 0$.

Full-compromise equilibrium: existence. In a full-compromise equilibrium, we have $c = 0$ and $\lambda = 1$. As $B \geq 0$, we have $w_1 \geq w_0$, and a necessary and sufficient condition for existence is then $w_1 \leq 0$. This can be written $w_1 = \frac{\Phi d_1 + \Delta}{1 + \Delta} \leq 0$, or equivalently $\frac{\Phi(B+D)}{2} + \Delta \leq 0$. Let $F_1(h) = \frac{\Phi(h)(B+D)}{2} + \Delta(h)$. To sum up, we have shown that a full-compromise equilibrium exists if and only if $F_1(h) \leq 0$.

Partial-compromise equilibrium: existence. Let us first consider the case $w_0 \leq 0$ and $w_1 > 0$.

As shown above, a partial-compromise equilibrium is characterized by the equation $1 - \frac{\Phi d_1}{2} = \lambda + \Delta \frac{\lambda^2}{2}$, or equivalently $1 - \frac{\Phi(B+D)}{4} = \lambda + \Delta \frac{\lambda^2}{2}$. Such an equilibrium exists if and only if $w_0 \leq 0$ and $\lambda < 1$. The lower bound of the agreement set is $-1 + w_1 = 1 - 2\lambda$, it follows that $w_1 = 2(1 - \lambda)$.

We may thus write:

$$\begin{aligned} w_0 = (w_0 + w_1) - w_1 &= \frac{\Phi(d_0 + d_1) + \Delta\lambda[(1 - c) + (1 + c)]}{1 + \Delta\lambda} - 2(1 - \lambda) \\ &= \frac{\Phi D + 2\Delta\lambda - 2(1 - \lambda)(1 + \Delta\lambda)}{1 + \Delta\lambda}. \end{aligned}$$

Hence, we can write :

$$\begin{aligned} w_0 \leq 0 &\Leftrightarrow \frac{\Phi D}{2} + \Delta\lambda \leq (1 - \lambda)(1 + \Delta\lambda) \Leftrightarrow \lambda + \Delta\lambda^2 \leq 1 - \frac{\Phi D}{2} \\ &\Leftrightarrow 2\left(\lambda + \Delta\frac{\lambda^2}{2}\right) - \lambda \leq 1 - \frac{\Phi D}{2} \\ &\Leftrightarrow 2\left(1 - \frac{\Phi(B+D)}{4}\right) - \lambda \leq 1 - \frac{\Phi D}{2} \\ &\Leftrightarrow 1 - \frac{\Phi B}{2} \leq \lambda. \end{aligned}$$

To conclude, such a partial-compromise equilibrium exists if and only if:

$$\exists \lambda \in \left[1 - \frac{\Phi B}{2}, 1\right), \quad 1 - \frac{\Phi(B+D)}{4} = \lambda + \Delta \frac{\lambda^2}{2}. \quad (10)$$

Note that $\lambda \mapsto \lambda + \Delta \frac{\lambda^2}{2}$ is increasing and continuous on $\left[1 - \frac{\Phi B}{2}, 1\right]$. By application of the intermediate value theorem, we obtain that condition (10) is equivalent to

$$\left(1 - \frac{\Phi B}{2}\right) + \frac{\Delta}{2} \left(1 - \frac{\Phi B}{2}\right)^2 \leq 1 - \frac{\Phi(B+D)}{4} < 1 + \frac{\Delta}{2}, \quad (11)$$

Thus, a partial-compromise equilibrium with $w_0 \leq 0$ and $w_1 > 0$ exists if and only if two conditions are jointly satisfied. The first condition is $\frac{\Phi(B+D)}{2} + \Delta > 0$, equivalent to $F_1(h) > 0$. The second condition can be written as $\frac{\Phi(B-D)}{2} \geq \Delta \left(1 - \frac{\Phi B}{2}\right)^2$, equivalent to $F_0(h) \geq 0$.

Consider now the case $w_0 > 0$ and $w_1 \leq 0$. We obtain as above (replacing d_1 by d_0) that

$1 - \frac{\Phi(D-B)}{4} = \lambda + \Delta \frac{\lambda^2}{2}$. Following the same steps as above, we also obtain that

$$w_1 \leq 0 \quad \Leftrightarrow \quad 2 \left(1 - \frac{\Phi(D-B)}{4} \right) - \lambda \leq 1 - \frac{\Phi D}{2} \quad \Leftrightarrow \quad 1 + \frac{\Phi B}{2} \leq \lambda.$$

This last condition is incompatible with $\lambda < 1$, which must hold at a partial-compromise equilibrium. Hence, such a partial-compromise equilibrium cannot exist.

4. Equilibrium uniqueness. As Φ is increasing in h and Δ is decreasing in h , we have that $F_0(h) = \frac{\Phi(h)(B-D)}{2} - \Delta(h) \left(1 - \frac{\Phi(h)B}{2} \right)^2$ is increasing in h , with $F_0(0) = -\Delta(0) = -\frac{\beta}{1-\beta} < 0$ and $F_0(1) = \Phi(1) \frac{B-D}{2} = \frac{B-D}{2} > 0$. As F_0 is also continuous, by application of the intermediate value theorem, there exists a unique value $h_0 \in (0, 1)$ for which $F_0(h_0) = 0$.

Similarly, the function $F_1(h) = \frac{\Phi(h)(B+D)}{2} + \Delta(h)$ is decreasing in h , with $F_1(0) = \Delta(0) = \frac{\beta}{1-\beta} > 0$ and $F_1(1) = \Phi(1) \frac{B+D}{2} = \frac{B+D}{2} < 0$. Hence, there exists a unique value $h_1 \in (0, 1)$ for which $F_1(h_1) = 0$.

Now, observe that:

$$\frac{\Phi}{\Delta}(h_1) = \frac{2}{-B-D} \geq \frac{2}{B-D} \geq \frac{2}{B-D} \left(1 - \frac{\Phi(h_0)B}{2} \right)^2 = \frac{\Phi}{\Delta}(h_0).$$

As $\frac{\Phi}{\Delta}(h)$ is decreasing on $(0, 1)$, we have $h_0 \leq h_1$, with an equality if and only if $B = 0$. To conclude, for any $h \in [0, 1]$, there exists a unique equilibrium:

- if $h < h_0$, there is a no-compromise equilibrium, and only this equilibrium, since $F_0(h) < 0$ and $F_1(h) > 0$.
- if $h_0 \leq h < h_1$, there is a partial-compromise equilibrium with $w_0 \leq 0$ and $w_1 > 0$, and only this equilibrium, since $F_0(h) \geq 0$ and $F_1(h) > 0$.
- if $h \geq h_1$, there is a full-compromise equilibrium, and only this equilibrium, since $F_0(h) < 0$ and $F_1(h) \leq 0$.

5. Comparative statics. Since Φ is increasing in h , it is immediate that $c = \frac{\Phi B}{2}$ is a non-decreasing function of h , strictly increasing whenever $B > 0$. From above, for h small enough, the

agreement probability λ satisfies $\Delta\lambda^2 + \lambda - 1 + \frac{\Phi D}{2} = 0$. Differentiating this equation, we obtain:

$$(1 + 2\lambda)\frac{\partial\lambda}{\partial h} = -\frac{D}{2}\frac{\partial\Phi}{\partial h} - \lambda^2\frac{\partial\Delta}{\partial h} > 0.$$

Hence, as $1 + 2\lambda > 0$, we have that λ is an increasing function of h . ■

A.1.2 Proof of **Theorem 2**

We start by proving the following lemma, which characterizes parties' joint welfare.

Lemma 1 *Parties' joint welfare at equilibrium solely depends on the agreement probability λ , and is given by:*

$$W = 1 - \lambda + \lambda^2.$$

Proof. Let w be an equilibrium and let A be its agreement set. We apply (1) and (3):

$$\begin{aligned} W_\theta &= \frac{(1 - \lambda) h d_\theta + \lambda \mathbb{E}[u_\theta(x) \mid x \in A]}{1 - \beta(1 - h)(1 - \lambda)}, \\ w_\theta &= \frac{h d_\theta + \beta(1 - h) \lambda \mathbb{E}[u_\theta(x) \mid x \in A]}{1 - \beta(1 - h)(1 - \lambda)}. \end{aligned}$$

Solving for $\lambda \mathbb{E}[u_\theta(x) \mid x \in A]$ in both expressions, we have:

$$\begin{aligned} \frac{\lambda \mathbb{E}[u_\theta(x) \mid x \in A]}{1 - \beta(1 - h)(1 - \lambda)} &= W_\theta - \frac{(1 - \lambda) h}{1 - \beta(1 - h)(1 - \lambda)} d_\theta \\ &= \frac{1}{\beta(1 - h)} \left(w_\theta - \frac{h}{1 - \beta(1 - h)(1 - \lambda)} d_\theta \right). \end{aligned}$$

Thus, simplifying, we have the affine relation:

$$W_\theta = \frac{1}{\beta(1 - h)} (w_\theta - h d_\theta). \quad (12)$$

Since w is an interior equilibrium, we have $w_0 + w_1 = 2(1 - \lambda)$, and we may write :

$$W = \frac{1 - \lambda - \frac{hD}{2}}{\beta(1 - h)}.$$

Applying (7), we obtain:

$$\begin{aligned}
 1 - \frac{\Phi D}{2} = \lambda + \Delta \lambda^2 &\Leftrightarrow 1 - \frac{h}{1 - \beta(1 - h)} \times \frac{D}{2} = \lambda + \frac{\beta(1 - h)}{1 - \beta(1 - h)} \lambda^2 \\
 &\Leftrightarrow 1 - \beta(1 - h) - \frac{hD}{2} = (1 - \beta(1 - h)) \lambda + \beta(1 - h) \lambda^2 \\
 &\Leftrightarrow 1 - \lambda - \frac{hD}{2} = \beta(1 - h) (1 - \lambda + \lambda^2).
 \end{aligned}$$

We thus obtain:

$$W = \frac{1 - \lambda - \frac{hD}{2}}{\beta(1 - h)} = 1 - \lambda + \lambda^2,$$

as desired. ■

We now prove **Theorem 2**.

Proof. By application of **Lemma 1**, we have $\frac{\partial W}{\partial h} = (2\lambda - 1) \frac{\partial \lambda}{\partial h}$. As $\frac{\partial \lambda}{\partial h}|_{h=0} > 0$, the effect of brinkmanship on welfare, that is $\frac{\partial W}{\partial h}|_{h=0}$, is of the same sign as $\lambda^0 - 1/2$. Applying (8), we obtain:

$$\lambda^0 := \lambda(h = 0) = \frac{1}{2\Delta^0}(\sqrt{1 + 4\Delta^0} - 1) = \frac{1 + 4\Delta^0 - 1}{2\Delta^0(\sqrt{1 + 4\Delta^0} + 1)} = \frac{2}{\sqrt{1 + 4\Delta^0} + 1}, \quad (13)$$

with $\Delta^0 = \Delta(h = 0)$. We then observe that (with the same equivalence between weak inequalities):

$$\lambda^0 > 1/2 \Leftrightarrow \Delta^0 = \frac{\beta}{1 - \beta} < 2 \Leftrightarrow \beta < 2/3. \quad (14)$$

We thus obtain $\frac{\partial W}{\partial h}|_{h=0} < 0$ when $\beta > 2/3$, while $\frac{\partial W}{\partial h}|_{h=0} > 0$ when $\beta < 2/3$.

To see how the effect of h on W is mediated by the crisis severity D , we may write $\frac{\partial W}{\partial h}|_{h=0} = (2\lambda^0 - 1) \frac{\partial \lambda}{\partial h}|_{h=0}$. As λ^0 does not depend on D , we look for an expression of $\frac{\partial \lambda}{\partial h}|_{h=0}$. Differentiating (7) with respect to h , we obtain:

$$0 = \frac{D}{2} \frac{\partial \Phi}{\partial h} + \lambda^2 \frac{\partial \Delta}{\partial h} + \frac{\partial \lambda}{\partial h} (1 + 2\Delta \lambda).$$

As $\frac{\partial \Phi}{\partial h}|_{h=0} = \frac{1}{1 - \beta}$ and $\frac{\partial \Delta}{\partial h}|_{h=0} = \frac{-\beta}{(1 - \beta)^2}$, we may write:

$$\frac{\partial \lambda}{\partial h}|_{h=0} = \frac{-\frac{D}{2(1 - \beta)} + \frac{\beta(\lambda^0)^2}{(1 - \beta)^2}}{1 + 2\Delta^0 \lambda^0}. \quad (15)$$

We see that $\frac{\partial \lambda}{\partial h}|_{h=0}$ is higher when the crisis is more severe (higher $|D|$). We thus obtain that the effect of h on W is stronger in that case. As both λ^0 and $\frac{\partial \lambda}{\partial h}|_{h=0}$ are independent of B , it follows that $\frac{\partial W}{\partial h}|_{h=0}$ does not depend on B . ■

A.1.3 Proof of Proposition 3

Proof. Applying equations (3), (4) and (5), we have:

$$\begin{aligned} W &= \frac{W_0 + W_1}{2} = \frac{1}{2} \left(\frac{(1-\lambda)hd_0 + \lambda(1-c)}{1-\beta(1-h)(1-\lambda)} + \frac{(1-\lambda)hd_1 + \lambda(1+c)}{1-\beta(1-h)(1-\lambda)} \right) \\ &= \frac{(1-\lambda)h\frac{D}{2} + \lambda}{1-\beta(1-h)(1-\lambda)} = \frac{\lambda + (1-\lambda)h\frac{D}{2}}{1-(1-h)(1-\lambda)} \times f(h), \end{aligned}$$

where $f(h) = \frac{1-(1-h)(1-\lambda)}{1-\beta(1-h)(1-\lambda)}$. Moreover, we may write:

$$\begin{aligned} \frac{W}{f(h)} &= \frac{\lambda + (1-\lambda)h\frac{D}{2}}{1-(1-h)(1-\lambda)} = 1 + \left(\frac{\lambda}{1-(1-h)(1-\lambda)} - 1 \right) + \frac{(1-\lambda)h\frac{D}{2}}{1-(1-h)(1-\lambda)} \\ &= 1 + \frac{\lambda - 1 + (1-h)(1-\lambda)}{1-(1-h)(1-\lambda)} + \frac{(1-\lambda)h\frac{D}{2}}{1-(1-h)(1-\lambda)} \\ &= 1 + \left(\frac{D}{2} - 1 \right) \frac{(1-\lambda)h}{1-(1-h)(1-\lambda)}. \end{aligned}$$

To conclude, observe first that:

$$\mathbb{P}(d^* | A) = (1-\lambda)h \sum_{k=0}^{+\infty} (1-\lambda)^k (1-h)^k = \frac{(1-\lambda)h}{1-(1-\lambda)(1-h)}.$$

Second, observe that $(1-h)(1-\lambda)$ decreases with h , while $\frac{1-x}{1-\beta x}$ decreases with x , as $\beta < 1$.

Together, this implies that the timing factor $f(h)$ is indeed increasing in h . ■

A.1.4 Proof of Proposition 4

Proof. Let w be an equilibrium and let A be its agreement set. We know from Lemma 1 that $W = 1 - \lambda + \lambda^2$. Since w is interior, we have $c = \frac{(-1+w_1) + (1-w_0)}{2}$, so that, using (6), we

obtain $w_1 - w_0 = 2c = \Phi B$. Using (12), we may write:

$$\begin{aligned} W_1 - W_0 &= \frac{1}{\beta(1-h)} (w_1 - w_0 - h(d_1 - d_0)) \\ &= \frac{1}{\beta(1-h)} (\Phi B - hB) = \frac{hB}{\beta(1-h)} \left(\frac{1}{1-\beta(1-h)} - 1 \right) = \Phi B. \end{aligned}$$

■

A.1.5 Proof of Theorem 3

Proof. 1. Formally, we show in this proof that: for any $B > 0$, there exists $\bar{\beta} \in (0, 1)$ such that for any $\beta \geq \bar{\beta}$, we have $\frac{\partial W_1}{\partial h}(h=0) > 0$.

By combining Lemma 1 and Proposition 4, we obtain

$$W_1 = \frac{W_0 + W_1}{2} + \frac{W_1 - W_0}{2} = 1 - \lambda + \lambda^2 + \frac{B\Phi}{2}.$$

We thus have $\frac{\partial W_1}{\partial h} = (2\lambda - 1)\frac{\partial \lambda}{\partial h} + \frac{B}{2}\frac{\partial \Phi}{\partial h}$. Let us study the two terms separately.

Letting $\delta = 1 - \beta$, we have that $\frac{\partial \Phi}{\partial h}|_{h=0} = \frac{1}{1-\beta} = \frac{1}{\delta}$.

Applying (15), we obtain:

$$\left. \frac{\partial \lambda}{\partial h} \right|_{h=0} = \frac{-\frac{D}{2\delta} + \frac{(1-\delta)(\lambda^0)^2}{\delta^2}}{1 + \frac{2\lambda^0(1-\delta)}{\delta}} = \frac{-\frac{D\delta}{2} + (1-\delta)(\lambda^0)^2}{\delta^2 + 2\lambda^0(1-\delta)\delta}. \quad (16)$$

Applying (13), we obtain, for small δ :

$$\lambda^0 = \frac{2}{\sqrt{1+4\Delta^0}+1} = \frac{2}{1+\sqrt{1+4\frac{1-\delta}{\delta}}} = \frac{2}{\sqrt{\frac{1}{\delta}}(\sqrt{\delta}+\sqrt{4-3\delta})} \sim \sqrt{\delta}.$$

We thus have, for small δ :

$$\left. \frac{\partial \lambda}{\partial h} \right|_{h=0} \sim \frac{-\frac{D\delta}{2} + (1-\delta)\delta}{\delta^2 + 2\sqrt{\delta}(1-\delta)\delta} \sim \frac{-\frac{D}{2} + (1-\delta)}{2\sqrt{\delta}(1-\delta)} \sim \frac{1 - \frac{D}{2}}{2\sqrt{\delta}}.$$

To sum up, the welfare W_1 of the advantaged party satisfies, for small δ , the following condition:

$$\left. \frac{\partial W_1}{\partial h} \right|_{h=0} = (2\lambda^0 - 1) \left. \frac{\partial \lambda}{\partial h} \right|_{h=0} + \frac{B}{2} \left. \frac{\partial \Phi}{\partial h} \right|_{h=0} \sim -\frac{1 - \frac{D}{2}}{2\sqrt{\delta}} + \frac{B}{2\delta} \sim \frac{B}{2\delta}.$$

Hence, for $B > 0$ we obtain that $\left. \frac{\partial W_1}{\partial h} \right|_{h=0} > 0$ for δ small enough, that is, for β close enough to 1.

2. When $D = -B$, party 1's utility derivative at $h = 0$ writes

$$\begin{aligned} \left. \frac{\partial W_1}{\partial h} \right|_{h=0} &= (2\lambda^0 - 1) \left. \frac{\partial \lambda}{\partial h} \right|_{h=0} + \frac{B}{2} \left. \frac{\partial \Phi}{\partial h} \right|_{h=0} \\ &= (2\lambda^0 - 1) \frac{\beta(\lambda^0)^2 + \frac{B(1-\beta)}{2}}{(1-\beta)^2 + 2\beta(1-\beta)\lambda^0} + \frac{B}{2(1-\beta)} \\ &= \frac{2(2\lambda^0 - 1)\beta(\lambda^0)^2 + B(1-\beta)(2\lambda^0 - 1) + B(1-\beta) + 2B\beta\lambda^0}{2(1-\beta)^2 + 4\beta(1-\beta)\lambda^0} \\ &= \frac{2(2\lambda^0 - 1)\beta(\lambda^0)^2 + 2B\lambda^0}{2(1-\beta)^2 + 4\beta(1-\beta)\lambda^0} \\ &= \frac{\lambda^0}{(1-\beta)^2 + 2\beta(1-\beta)\lambda^0} (\beta\lambda^0(2\lambda^0 - 1) + B). \end{aligned}$$

Hence:

- for any $\beta < 2/3$, we have $\lambda^0 < 1/2$ and thus $\left. \frac{\partial W_1}{\partial h} \right|_{h=0} > 0$.
- for any $\beta \geq 2/3$, we have $\left. \frac{\partial W_1}{\partial h} \right|_{h=0} > 0$ whenever $B > \beta\lambda^0|2\lambda^0 - 1|$.

Hence, a necessary and sufficient condition to have $\left. \frac{\partial W_1}{\partial h} \right|_{h=0} > 0$ when $D = -B$ for any β is

$$B > \min_{\beta \geq 2/3} \beta\lambda^0|2\lambda^0 - 1| := \bar{B}$$

As $\lambda|2\lambda - 1| \leq \frac{1}{8}$, we obtain that $\bar{B} < \frac{1}{8}$.¹⁸

To conclude, we have shown that for any $D < -\frac{1}{8}$, party 1 has an incentive to activate the brinkmanship strategy if she is advantaged enough (if $B = -D$, and by continuity, if $B < -D$ is high enough), independently on the discount factor β . ■

¹⁸We use this upper bound on $\bar{B}$ for simplicity in the statement of the result. Numerically, we find that $\bar{B} \approx 0.116$.

A.1.6 Proof of Proposition 5

Proof. As $W_0 = 1 - \lambda + \lambda^2 - \frac{B\Phi}{2}$, we obtain by application of (16):

$$\begin{aligned} \left. \frac{\partial W_0}{\partial h} \right|_{h=0} &= (2\lambda^0 - 1) \left. \frac{\partial \lambda}{\partial h} \right|_{h=0} - \frac{B}{2} \left. \frac{\partial \Phi}{\partial h} \right|_{h=0} = (2\lambda^0 - 1) \frac{\beta(\lambda^0)^2 - \frac{D(1-\beta)}{2}}{(1-\beta)^2 + 2\beta(1-\beta)\lambda^0} - \frac{B}{2(1-\beta)} \\ &\geq (2\lambda^0 - 1) \frac{\beta(\lambda^0)^2 + \frac{B(1-\beta)}{2}}{(1-\beta)^2 + 2\beta(1-\beta)\lambda^0} - \frac{B}{2(1-\beta)}. \end{aligned}$$

By application of (14), we know that $\lambda^0 > 1/2$ when $\beta < 2/3$. In that case, we thus know that the previous expression is positive for $B = 0$ (and thus by continuity for B small enough), hence the result. ■

A.2 Intermediate and large brinkmanship threats

While the core of the paper focused on the most relevant case of a small brinkmanship threat, we now discuss how the model changes with a larger threat. We start by characterizing the (unique) stationary equilibrium for all possible brinkmanship threats, thus generalizing Theorem 1.

Theorem 4 *For any $h \in [0, 1]$, there exists a unique stationary equilibrium. There exist thresholds $h_0, h_1 \in (0, 1)$ with $h_0 \leq h_1$ such that:*

- *if $h < h_0$, the equilibrium is interior: both parties reject some deals,*
- *if $h_0 \leq h < h_1$, the equilibrium exhibits partial compromise: party 0 accepts all deals, while party 1 rejects some of them,*
- *if $h \geq h_1$, the equilibrium exhibits full compromise: both parties accept all deals.*

As the proof of Theorem 1 in Appendix A.1 exhaustively covers all levels of threats, it actually proves Theorem 4. To illustrate the result, we display the bounds of the equilibrium agreement set on Figure 6 (which extends Figure 3), for a numerical example.

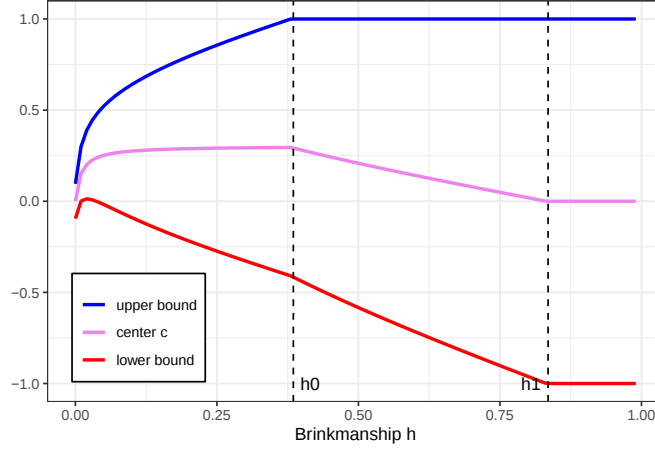


Figure 6: Agreement set for $\beta = 0.99$, $D = -1$ and $B = 0.6$.

We now discuss the regimes with intermediate and large brinkmanship threats, and conclude with a welfare analysis.

Intermediate brinkmanship threat: $h_0 \leq h < h_1$. In this region, we obtained in the proof of [Theorem 4](#) that $\lambda = \frac{1}{\Delta} \left(\sqrt{1 + 2\Delta \left(1 - \frac{\Phi d_1}{2}\right)} - 1 \right)$ and $c = 1 - \lambda$. As for small h , the agreement probability λ is increasing with h . Intuitively, a higher threat forces party 1 to compromise more, while party 0 always compromises. As a result, the center of the agreement set c decreases with the threat h in that region.

Large brinkmanship threat: $h \geq h_1$. In this region, both parties always compromise. We thus have $\lambda = 1$ and $c = 0$. Intuitively, the threat is so high that the first deal ever proposed is accepted by both parties.

Welfare. We illustrate parties' welfare on [Figure 7](#) (which extends [Figure 4](#)), for a numerical example.

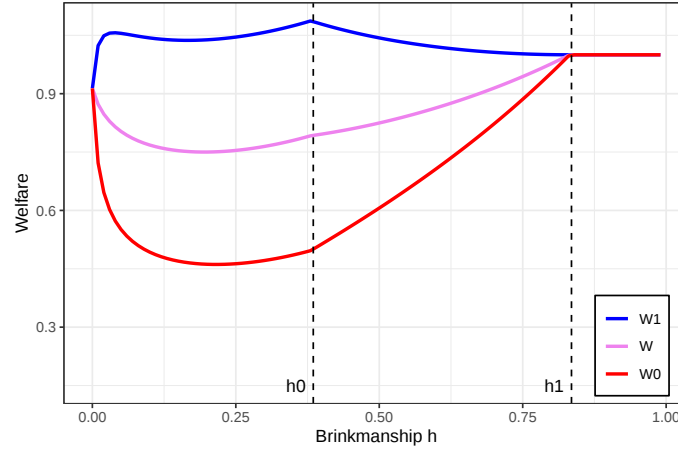


Figure 7: Parties' welfare for $\beta = 0.99$, $D = -1$ and $B = 0.6$

Figure 7 illustrates the following general results. Overall, parties' joint welfare is maximal for a large threat ($h \geq h_1$), since a deal is accepted without any delay. For intermediate threats, party 0's welfare is always lower than with a high threat, for two reasons: delay occurs on the equilibrium path, and accepted deals are on average more favorable to party 1. Meanwhile, one can show that party 1's welfare decreases with h within that regime.¹⁹ Hence, if the advantaged party could choose the threat level, it would pick it optimally in the interval $(0, h_0]$.

References

- Avidit Acharya and Edoardo Grillo. War with crazy types. *Political Science Research and Methods*, 3(2):281–307, 2015.
- Avidit Acharya and Juan Ortner. Paths to the frontier. *American Economic Journal: Microeconomics*, 14(1):39–69, 2022.
- James Albrecht, Axel Anderson, and Susan Vroman. Search by committee. *Journal of Economic Theory*, 145(4):1386–1407, 2010.
- Deepal Basak and Joyee Deb. Gambling over public opinion. *American Economic Review*, 110(11):3492–3521, 2020.

¹⁹The formula $W_1 = 2(1 - \lambda + \lambda^2/2)$ can be obtained by following the same steps as in the proof of Lemma 1 in Appendix A.1.

- Ken Binmore, Ariel Rubinstein, and Asher Wolinsky. The nash bargaining solution in economic modelling. *The RAND Journal of Economics*, pages 176–188, 1986.
- T. Renee Bowen, Ying Chen, Hülya Eraslan, and Jan Zápál. Efficiency of flexible budgetary institutions. *Journal of Economic Theory*, 167:148–176, 2017.
- Shirish D. Chikte and Sudhakar D. Deshmukh. The role of external search in bilateral bargaining. *Operations Research*, 35(2):198–205, 1987.
- Olivier Compte and Philippe Jehiel. Bargaining and majority rules: A collective search perspective. *Journal of Political Economy*, 118(2):189–221, 2010.
- Peter C. Cramton and Joseph S. Tracy. Strikes and holdouts in wage bargaining: theory and data. *American Economic Review*, 82:100–121, 1992.
- Pablo Cuellar and Lucas Rentschler. Threat, commitment, and brinkmanship in adversarial bargaining. Technical report, 2023.
- Bhaskar Dutta and Arunava Sen. Nash implementation with partially honest individuals. *Games and Economic Behavior*, 74(1):154–169, 2012.
- Wioletta Dziuda and Antoine Loeper. Dynamic collective choice with endogenous status quo. *Journal of Political Economy*, 124(4):1148–1186, 2016.
- Tore Ellingsen and Topi Miettinen. Commitment and conflict in bilateral bargaining. *American Economic Review*, 98(4):1629–35, 2008.
- Hülya Eraslan and Antonio Merlo. Majority rule in a stochastic model of bargaining. *Journal of Economic Theory*, 103(1):31–48, 2002.
- William Fuchs and Andrzej Skrzypacz. Bargaining with deadlines and private information. *American Economic Journal: Microeconomics*, 5(4):219–43, 2013.
- Timothy F. Geithner. Secretary geithner sends debt limit letter to congress (1/6/2011). *U.S. Department of the Treasury (available at <https://home.treasury.gov/secretary-geithner-sends-debt-limit-letter-to-congress-2>)*, 2011.

- Uri Gneezy, Ernan Haruvy, and Alvin E Roth. Bargaining under a deadline: Evidence from the reverse ultimatum game. *Games and Economic Behavior*, 45(2):347–368, 2003.
- Edoardo Grillo and Carlo Prato. Reference points and democratic backsliding. *American Journal of Political Science*, 2020.
- Albert Ma and Michael Manove. Bargaining with deadlines and imperfect player control. *Econometrica*, 61(6):1313–1339, 1993.
- Paola Manzini and Marco Mariotti. Going alone together: joint outside options in bilateral negotiations. *The Economic Journal*, 114(498):943–960, 2004.
- Antonio Merlo and Charles Wilson. A stochastic model of sequential bargaining with complete information. *Econometrica*, 63(2):371–399, 1995.
- Antonio Merlo and Charles Wilson. Efficient delays in a stochastic model of bargaining. *Economic Theory*, 11(1):39–55, 1998.
- Matthias Messner and Mattias K. Polborn. The option to wait in collective decisions and optimal majority rules. *Journal of Public Economics*, 96(5-6):524–540, 2012.
- Benny Moldovanu and Frank Rosar. Brexit: A comparison of dynamic voting games with irreversible options. *Games and Economic Behavior*, 130:85–108, 2021.
- Benny Moldovanu and Xianwen Shi. Specialization and partisanship in committee search. *Theoretical Economics*, 8(3):751–774, 2013.
- Abhinay Muthoo. On the strategic role of outside options in bilateral bargaining. *Operations Research*, 43(2):292–297, 1995.
- Juan Ortner. A theory of political gridlock. *Theoretical Economics*, 12(2):555–586, 2017.
- John W. Patty. Signaling through obstruction. *American Journal of Political Science*, 60(1):175–189, 2016.
- Elizabeth Maggie Penn. A model of farsighted voting. *American Journal of Political Science*, 53(1):36–54, 2009.

- Robert Powell. Nuclear brinkmanship with two-sided incomplete information. *American Political Science Review*, 82(1):155–178, 1988.
- Michael Schwarz and Konstantin Sonin. A theory of brinkmanship, conflicts, and commitments. *The Journal of Law, Economics, & Organization*, 24(1):163–183, 2008.
- Alp Simsek and Muhamet Yildiz. Durability, deadline, and election effects in bargaining. Technical report, National Bureau of Economic Research, 2016.
- Bruno Strulovici. Learning while voting: Determinants of collective experimentation. *Econometrica*, 78(3):933–971, 2010.
- Asher Wolinsky. Matching, search, and bargaining. *Journal of Economic Theory*, 42(2):311–333, 1987.