

Modern Force and Coercive Backfire

Shusuke Ioku[†]

September 21, 2026

ABSTRACT. Modern military technologies, including automated weapons and AI-assisted systems, are thought to strengthen deterrence and preserve peace by enhancing limited military operations without provoking wider war. I challenge this optimism by examining their implications for coercive bargaining rather than their direct military effects. I develop a coercive bargaining model with a *subwar capability*: access to limited operations short of general war. Such operations can replace war after resistance or settlement after compliance, potentially weakening coercive threats and assurances. The target may concede less, further weakening assurances and the government's ability to secure concessions. Across broad parameter regions, I show that subwar capability can lower coercive ability and increase war, even when no operation is used. Subwar capability can weaken the effect of military power on coercive ability and reduce social welfare by raising conflict costs. More efficient military technology can thus disrupt coercive bargaining and raise the cost of conflict.

KEYWORDS. coercive bargaining; military technology; limited military operations; assurance; formal theory

Word Count: 9,424

[†]Ph.D. Candidate, Department of Political Science, University of Rochester.
Email: iokushusuke@gmail.com.

INTRODUCTION

Violence has declined over the long run, argues [Pinker \(2011\)](#). This trend seems particularly evident in interstate politics, where states increasingly pursue limited gains while seeking to avoid war ([Gat 2013](#); [Gleditsch and Pickering 2013](#); [Altman 2020](#)). Automated weapons and AI-assisted military systems promise to advance this trend by strengthening deterrence and preserving peace ([Zegart 2018](#); [Jensen, Whyte, and Cuomo 2019](#); [Horowitz 2020](#)). Kathleen Hicks, then U.S. deputy defense secretary, argued that AI-enabled decisions could be “decisive in deterring a fight,” while Ukraine’s defense AI chief Danylo Tsvok said democracies cannot “effectively protect peace” without it.^[1] A key role practitioners emphasize is enabling governments to pursue limited objectives through military operations without provoking wider war. Retired Indian general A. B. Shivane links autonomy and precision strikes to “swift, limited, conventional action without breaching escalation thresholds” and credible deterrence, with AI-assisted air defence helping control escalation.^[2] I call access to limited operations short of general war a *subwar capability*. The appeal is that such operations could reduce the costs of interstate conflict by substituting for general war.

I challenge this optimism about modern military technology by examining its implications for coercive bargaining rather than its direct military effects. I show that subwar capability can weaken a government’s ability to secure concessions without war and raise the overall cost of interstate conflict. The subwar option can change what the target expects from resisting or complying, and therefore how much it is willing to concede. Coercion requires both a credible *threat* to punish resistance and a credible *assurance* that compliance will prevent further harm ([Schelling 1960](#)). The subwar option can weaken threats when it replaces a more damaging war and assurances when it replaces settlement. Either change reduces what the target will concede. Smaller concessions reduce the demanding

^[1]Kathleen Hicks, “[AI in the Era of Strategic Competition](#),” March 5, 2024; Derek Gatopoulos and Hanna Arhirova, [interview with Danylo Tsvok](#), Associated Press, April 2026.

^[2]A. B. Shivane, [CS Conversations](#), October 10, 2025.

government's incentive for restraint, further weakening assurances, and can make general war more attractive, even when the subwar operation is never used.

I develop this argument in a parsimonious framework of coercive bargaining, in which one state, the coercer, demands a concession, the target complies or resists, and the coercer then chooses whether to use force (Banks 1990; Fearon 1994; Schultz 1999). The target does not know how costly the coercer would find a war. I compare a benchmark with settlement and general war to a game that adds the subwar option after either response. Following the experimental approach to formal modeling in Paine and Tyson (2020), I vary only the availability of the subwar option, holding all other features of the bargaining environment fixed. In each game, I select the equilibrium with the largest accepted demand. I measure *coercive ability* as the expected concession retained without general war: what the target gives up, discounted by the risk that war reopens the dispute. I also assess the coercer's total payoff, the probability of war, and the two states' combined conflict costs.

This bargaining mechanism can weaken the translation of military power into coercive ability. Without the subwar option, greater military power makes resistance more dangerous, so the target is willing to concede more to avoid war. That larger concession also gives the coercer more to preserve through settlement, reducing its incentive to fight. A subwar operation can disrupt this exchange. The coercer keeps the target's concession even if it conducts the operation afterward. An extra concession therefore raises the payoff from stopping and from operating by the same amount; it does not make stopping more attractive relative to the operation. The target consequently cannot count on a larger concession to remove this source of further harm. Its willingness to concede depends on how much safer compliance makes it, as well as how dangerous resistance becomes. When the subwar option erodes that protection, the target may respond less to an increase in military power. The resulting smaller concession gives the coercer less to preserve through restraint and can make war more attractive after compliance. Greater military power can therefore produce a smaller increase in the concession the coercer expects to retain without war.

The resulting reduction in concessions can also increase the risk of general war, despite the availability of a limited alternative. With the target making the same concession or resisting as before, the coercer may choose a subwar operation where it would otherwise wage general war. Anticipating this less damaging response to resistance, the target may be willing to concede less. The prospect of a subwar operation even after compliance can further reduce its willingness to concede. A smaller concession leaves the coercer less to preserve through restraint, making general war more attractive. If the concession falls sufficiently, this bargaining response outweighs the direct substitution away from war. The subwar option can have the immediate effect of replacing war with limited force, yet the ultimate effect of making war more likely by undermining coercive threats and assurances, thereby inducing smaller concessions.

These changes in bargaining and the use of force can also raise the two states' combined conflict costs, even as the cost of conducting a subwar operation falls. Replacing war with a subwar operation saves costs, while replacing settlement adds them. When the added costs exceed the savings, a reduction in operating cost increases expected conflict costs. When lower operating cost also reduces the target's concession, that bargaining response makes war more attractive and adds further costs. Evaluating a capability only by the expense of each operation or the number of general wars can therefore miss the costs it creates by changing bargaining and the use of force.

The adverse outcomes arise together over broad parameter regions. For every military and technological configuration permitting the active, interior comparison, I establish a common interval of backing-down costs over which coercive ability and the coercer's expected payoff fall, military power yields smaller bargaining gains, and war and conflict costs rise. Several comparisons hold beyond the closed-form restrictions. On this common interval, the coercive and payoff losses and the increases in war and conflict cost also persist under small changes in war-cost uncertainty and escalation risk.

These findings contribute to research on the benefits and risks of force short of war. Limited seizures can secure gains without provoking war ([Tarar 2016](#); [Altman 2017](#); [2020](#);

Maass 2021), while low-level operations can avert preventive war by slowing an adversary's rise (Schram 2021). Other work identifies risks: less destructive retaliation can make nuclear exchanges more likely (Powell 1989), and improved capabilities to counter low-level aggression can induce riskier challenges and war (Gannon et al. 2024). Schram (2022) shows that better low-level capabilities can leave every type of their holder no better off through greater predictability or increased risk taking by its opponent. Spaniel and İdrisoğlu (2024) show that switching to a strategy with lower costs of fighting can reduce demands and increase war risk; with shifting power, both states can lose. In their incomplete-information models, the target privately knows its military costs or strength. Here, the coercer privately knows its war cost. This difference facilitates analysis of how subwar capability affects coercive ability through the target's assessment of threats and assurances.

The paper also speaks to research on threats and assurances in coercive bargaining. Coercion requires both costly resistance and protection through compliance (Schelling 1966; Cebul, Dafoe, and Monteiro 2021; Pauly 2024; 2025). Pape (1996) classifies strategies by how they make resistance costly or futile. Kydd and McManus (2017) show that threats and assurances can both be credible and that anticipated attack after a concession can reduce what the target will concede. Ioku and Kurizaki (2026) show how the prospect of attack after compliance can allow coercers planning war to imitate the demands of those seeking settlement. Adding the same subwar option after either response can connect threats and assurances through the concession. Even when the operation only weakens punishment for resistance, the target can concede less, making war more attractive after compliance. I thus identify a channel through which weaker threats undermine assurance by reducing what the coercer has to preserve through restraint.

The paper also contributes to research on the relationship between military power and coercive ability. The capacity to impose costs need not track the capacity to win a war (Slantchev 2003; 2011; Kenkel and Schram 2025). Sechser (2010) shows that weaker targets may resist militarily superior coercers because compliance can reveal low resolve and invite further demands. More generally, powerful coercers can fail when they cannot

credibly promise restraint (Art and Greenhill 2018; Cebul, Dafoe, and Monteiro 2021). Here, the coercer retains the concession even if it conducts a subwar operation, so a larger concession does not make settlement more attractive relative to operating. I show that this assurance problem can make an increase in military power produce a smaller gain in expected concessions retained without war when the subwar option is available.

THE FRAMEWORK

Figure 1 shows the timing of the coercive bargaining game. The coercer makes a take-it-or-leave-it demand, the target concedes or resists, and the coercer chooses a continuation action. The coercer's cost of war is private information, and abandoning a resisted demand carries an audience cost (Banks 1990; Fearon 1994; Schultz 1999). The benchmark permits settlement or war; the second game adds a subwar operation after either target response.

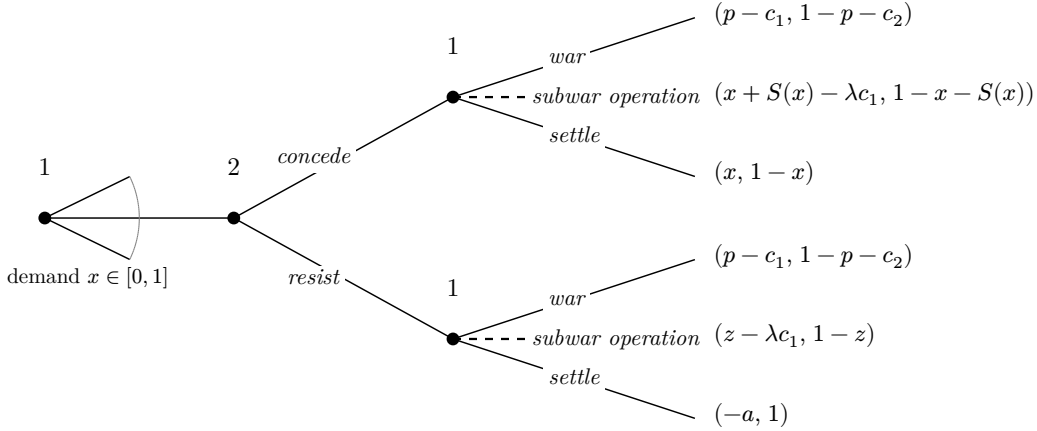


Figure 1: Coercive Bargaining Game with and without the Subwar Option

Note: The tree is conditional on the coercer's privately observed cost c_1 . Player 1 is the coercer and player 2 the target, which observes the demand but knows only the cost distribution; the coercer's payoff appears first. The dashed subwar-operation branches are added by the subwar capability and absent from the benchmark. War reopens the dispute, so its payoffs do not depend on the transfer. A subwar operation preserves the transfer, produces gain $S(x) = \min\{z, 1 - x\}$ after compliance and z after resistance, and costs λc_1 . Settling after resistance costs the coercer a .

Players, Actions, and Timing

Two states value a disputed good at 1. The *target*, player 2, initially holds the good. The *coercer*, player 1, has a war cost c_1 drawn uniformly from $[0, \bar{c}]$. The coercer observes its cost, but the target knows only its distribution F , with density $f = \frac{1}{\bar{c}}$ on its support. Throughout, $F(t) = 0$ for $t \leq 0$ and $F(t) = 1$ for $t \geq \bar{c}$. A lower cost denotes a more resolute coercer. The target's war cost $c_2 \geq 0$ is common knowledge.

The coercer demands $x \in [0, 1]$. The target chooses $r \in \{0, 1\}$, where $r = 1$ denotes resistance and $r = 0$ compliance. Let y be the transfer in place at the continuation stage. Thus $y = x$ after compliance and $y = 0$ after resistance. The coercer chooses from the same continuation menu after either response. A continuation history is (y, r) ; compliance with a zero demand is $(0, 0)$.

Settlement leaves the transfer y in place. If the target resisted, settlement also imposes an audience cost $a > 0$ on the coercer for abandoning its demand. War gives payoffs $(p - c_1, 1 - p - c_2)$, where $p \in (0, 1)$ is the coercer's probability of victory. Because war reopens the whole dispute, these payoffs do not depend on any earlier transfer. At continuation indifference the coercer chooses war before the operation and the operation before settlement.

The second game adds a *subwar operation* with reach $z > 0$. The baseline puts different operations on a common payoff scale. Let

$$S(y) = \min\{z, 1 - y\} \tag{1}$$

denote the resulting gain to the coercer and harm to the target. The operation costs the coercer λc_1 , where $\lambda \in (0, 1)$, and produces payoffs $(y + S(y) - \lambda c_1, 1 - y - S(y))$. The parameter measures the cost of conducting an operation relative to war, conditional on possessing the capability. Development and acquisition costs are outside the model.

The analysis maintains the following parameter region.

ASSUMPTION 1 (Parameter Restrictions): (i) The operation has limited reach,

$$p + c_2 + 2z \leq 1$$

and (ii) uncertainty about the coercer's war cost is sufficiently large,

$$\frac{2p + c_2}{1 - \lambda} < \bar{c}$$

A war costs the target $p + c_2$ relative to keeping the disputed good. [Assumption 1\(i\)](#) ensures that the operation has its full effect at every accepted demand, so $S(y) = z$ there.^[3] [Assumption 1\(ii\)](#) implies $\bar{c} > 2p + c_2$ in the benchmark and directly imposes the stronger bound needed in the operation game. These inequalities make the target's payoff from concession decrease with the demand under the prior. Additional conditions governing interior cutoffs and the coercer's continuation menu are stated with the results that use them.

Comments on the Model

The two games differ only in whether the coercer has a subwar option; the political environment and information are held fixed. Following [Paine and Tyson \(2020\)](#), I interpret this comparison as an ideal experiment that isolates how the capability changes threats and assurances. Both games allow force after compliance, so the subwar option changes incentives within a shared commitment problem. In each game, I select the equilibrium in which the coercer makes the largest demand the target accepts. [§ Scope of the Equilibrium Comparison](#) explains which conclusions depend on that selection.

Residual objectives. A subwar operation preserves any concession already made and pursues a residual objective that yields a fixed additional gain at accepted demands.

^[3]If $x > 1 - z$, the operation pays $1 - \lambda c_1 > p - c_1$ for every type, so no war follows compliance. The target then loses at least x , whereas its loss after resistance is at most $\max\{p + c_2, z\} \leq 1 - z$ under every posterior. Such a demand is therefore strictly rejected. This argument does not depend on pooling or equilibrium selection.

This reflects the capacity for limited action “without breaching escalation thresholds” emphasized by Shivane in the introduction. In the model, the coercer can pursue a limited objective without reopening the whole dispute or surrendering gains already secured through compliance. India presented the opening strikes of Operation Sindoor (May 2025) as “focused, measured and non-escalatory,” excluding Pakistani military facilities from the targets. During the operation, Akashteer combined radar data to automate air-defence coordination and responses to incoming threats. Operation Midnight Hammer (June 2025) used Tomahawk missiles with automated navigation to strike Iranian nuclear infrastructure; U.S. officials described its objective as disabling the nuclear program rather than regime change.^[4] The fixed incremental gain represents a residual objective whose value is unchanged across the concessions considered. It does not apply when compliance itself eliminates that objective.

Operating cost. Every coercer type pays the same fixed fraction of its war cost to conduct a subwar operation. This preserves the same ranking of coercer types across the two military actions. A reduction in λ lowers the cost of conducting a given operation while holding its effect z and the probability of victory p fixed.

Private costs. Coercer war costs are uniformly distributed. Uniformity delivers closed-form demands; the payoff-loss threshold and threat–assurance identity do not require it. Strict equilibrium comparisons also persist under small smooth departures from the uniform prior, as SI §B.5. [Robustness to Nearby Priors](#) establishes.

Gain and harm. The operation’s gain to the coercer equals its harm to the target, with the coercer’s operating cost deducted separately. Using one parameter for both keeps the baseline focused on the bargaining consequences of the additional military option. SI

^[4]See India’s Ministry of Defence, [statement of 7 May 2025](#) the Press Information Bureau, [account of Akashteer, 16 May 2025](#) the U.S. Department of Defense, [press briefing of 22 June 2025](#) and the U.S. Navy, [Tomahawk fact file](#), on onboard navigation and guidance. These cases illustrate limited stated objectives; the retention of earlier concessions is a modeling assumption.

[Corollary 1](#) allows gain and harm to differ and distinguishes their roles in the coercer's choice of force and the target's willingness to concede.

Power shift. The baseline assumes that war reopens the entire dispute and that prior concessions do not change the probability of victory. These assumptions keep war payoffs independent of the concession, allowing me to isolate how concessions change the coercer's incentive for restraint. [SI §F.2. Power Shift and Partial Recontestation](#) allows concessions to improve military prospects or remain partly secure during war and derives the conditions under which the core results on coercive ability and conflict continue to hold.

Escalation. The baseline excludes escalation from a subwar operation to war after the operation begins. Allowing an operation to escalate would add a direct route from limited force to war. Excluding that route isolates how the option's availability changes threats and assurances, the target's concession, and the coercer's incentive to choose war. [SI §F.1. Escalation Risk](#) relaxes the restriction and establishes that the core results on coercive ability and conflict remain valid for sufficiently small escalation risk under the conditions specified there.

Signaling. I omit a separate stage for signaling resolve through tying hands, sinking costs, or military threats ([Fearon 1997](#); [Slantchev 2005](#)). This avoids the additional complexity of analyzing signal choice and belief updating, allowing me to isolate how the subwar option changes threats, assurances, and concessions.

Solution Concept

I solve for pure-strategy perfect Bayesian equilibria satisfying the intuitive criterion ([Cho and Kreps 1987](#)). The target responds to each demand by either complying or resisting, without randomizing, and concedes when indifferent. After an unused demand, the criterion excludes a coercer type if every sequentially rational target response gives it strictly less than its equilibrium payoff, provided some type remains. Types that can match their equilibrium payoff remain admissible. [SI §C.1. Off-Path Beliefs and the Intuitive Criterion](#)

constructs supporting beliefs. Unless a result explicitly considers the broader class of perfect Bayesian equilibria, all equilibrium claims impose the criterion.

Equilibrium multiplicity makes a full characterization of demand strategies challenging. I focus on pooling equilibria, which reproduce the payoffs and military outcomes of every equilibrium with pure target responses, including any separating or partially separating equilibrium. Any coercer type that strictly benefits from compliance chooses the largest demand the target accepts. Only types indifferent between compliance and resistance can make other demands, and their different signals leave military actions and payoffs unchanged. SI [Lemma 5](#) proves that each such equilibrium has a pooling counterpart with the same military action and both states' payoffs for every coercer type.^[5]

I select the largest demand accepted under the prior in each game because it gives every coercer type its highest equilibrium payoff and maximizes the unconditional retained concession. It also minimizes unconditional war and expected joint conflict cost within the pure-response class, as SI [Corollary 2](#) shows: a larger concession makes settlement and the operation more attractive relative to war. An increase in war or conflict cost in the selected comparison therefore raises the respective minimum attainable across all equilibria in this class, including separating equilibria. The intuitive criterion does not uniquely select this equilibrium; § [Scope of the Equilibrium Comparison](#) explains the remaining selection and refinement limits.

THE BENCHMARK

Begin with the benchmark: set $z = 0$ and exclude the subwar operation. The coercer can choose only settlement or war after either target response. I compare the games in terms of the coercer's *coercive ability*, defined as follows.

^[5]This equivalence concerns payoffs and unconditional outcomes. Different signals can change outcomes conditional on compliance; the reported coercive-ability measure uses the pooling assessment.

DEFINITION 1 (Coercive Ability): A coercer’s *coercive ability* is the expected concession retained without general war:

$$\Omega = x [1 - \Pr(\text{war})]$$

where x is the conceded demand and $\Pr(\text{war})$ is the probability of war after compliance.

This measure excludes additional gains from the operation and continuation costs. Proposition 5 reports the coercer’s expected total payoff. I compare the largest accepted demands under the common selection stated above. § [Scope of the Equilibrium Comparison](#) discusses separating demands, alternative equilibria, and refinement limits. A dagger marks benchmark quantities and a star marks quantities in the game with the operation. The Supporting Information contains all proofs.

At a fixed demand, the operation cannot mechanically reduce coercive ability. If it replaces war after compliance, the transfer survives more often. If it replaces settlement, the transfer remains in place. Any reduction in coercive ability must therefore arise because the subwar capability changes the demand the target accepts.

To solve the benchmark, start with the coercer’s choice after the target responds. War becomes less attractive as its cost rises. After compliance, the coercer fights when its war cost is at or below $\kappa^\dagger(x) = p - x$ and settles otherwise. A larger concession lowers this cutoff because settlement preserves the transfer whereas war does not. After resistance, backing down instead carries the audience cost, raising the war cutoff to $p + a$. The full rule for either response is given in SI [Lemma 1](#).

Concession weakly raises every coercer type’s continuation value. A type that settles receives x after concession instead of paying a after resistance. A type that fights receives $p - c_1$ after either response and is indifferent. Thus no type prefers resistance, and every settling type prefers the largest accepted demand.

I therefore consider a pooling equilibrium in which the demand conveys no information about c_1 . Indifferent fighting types can make other demands without changing the acceptance condition at the largest accepted demand. Their signals preserve payoffs and unconditional military outcomes, although measures conditional on compliance can change. The reported coercive ability is for the pooling equilibrium.

The target's acceptance decision determines the pooling demand. Concession prevents war among types whose war costs fall between the two cutoffs. The target weighs this protection against the concession it gives up when war is avoided. It accepts as long as the expected war loss prevented is at least as large as that expected concession. At the largest accepted demand, $x^\dagger$, the target is indifferent.

PROPOSITION 1 (Pooling Sustains the Largest Benchmark Concession): Under [Assumption 1](#) and interior war cutoffs at every accepted demand, pooling at $x^\dagger$ satisfies the intuitive criterion and attains the largest demand sustainable under pure target responses. At that demand, type payoffs and unconditional military outcomes are unique. After compliance, the coercer fights when $c_1 \leq \kappa^\dagger(x^\dagger)$, so settlement has probability $\Pr^\dagger(\text{settlement} \mid x^\dagger) = 1 - F(\kappa^\dagger(x^\dagger))$. Coercive ability is given by [\(2\)](#).

The acceptance inequality and the formula for $x^\dagger$ are derived in [SI §C.3. Proof of Propositions 1 and 2](#).

Proposition 1 determines both the demand and continuation behavior. Types below the cutoff fight after concession, while types above it settle. Holding the demand fixed, greater military power raises the cutoff and makes war optimal for a larger set of coercer types. The equilibrium demand also rises with power, so the equilibrium war probability need not rise.

Rearranging the target's binding acceptance condition at $x^\dagger$ and substituting Definition 1 gives the subtraction form

$$\Omega^\dagger = \underbrace{(p + c_2) F(p + a)}_{\text{punishment after resistance}} - \underbrace{(p + c_2) F(\kappa^\dagger(x^\dagger))}_{\text{harm not prevented by compliance}} \quad (2)$$

Equation (2) writes coercive ability as one difference. The first term is the expected punishment produced by resistance. The second is the expected harm that compliance fails to prevent. A larger concession lowers the post-compliance war cutoff, makes settlement more likely, and reduces the second term. Greater power raises the first term and raises coercive ability as a whole, so whatever protection compliance loses is outweighed. The benchmark thus reproduces the expectation that power strengthens the threat faster than it erodes assurance (Cebul, Dafoe, and Monteiro 2021). The next section asks whether the subwar capability behaves the same way.

ADDING SUBWAR CAPABILITY

Now set $z > 0$ and add the subwar option. War costs rise faster with c_1 than the cost of a subwar operation, while settlement involves neither cost. When all three actions are optimal for some types, low-cost types fight, intermediate types use the operation, and high-cost types settle. Write $\kappa_W(y)$ for the boundary between war and the operation, and $\kappa_L(y, r)$ for the boundary between the operation and settlement. If these boundaries are reversed, the operation is dominated and the coercer chooses between war and settlement. SI Lemma 3 gives the cutoff formulas.

Write $\Pr^*(\text{war} \mid x)$ for the probability of war after compliance ($r = 0$) with demand x , and $\Pr^*(\text{war} \mid r = 1)$ for the probability after resistance. The same notation applies to the operation and settlement. Probabilities use the prior distribution of the coercer's war cost, or the posterior if the demand changes beliefs.

Before solving for the demand, hold transfer y fixed and compare threat and assurance. This isolates the operation's effect on the coercer's use of force before accounting for changes in the concession the target accepts. Some types that would fight without the option instead use the operation. When war is more harmful, this substitution weakens

the threat after resistance and strengthens assurance after compliance. Other types replace settlement with the operation. That substitution strengthens the threat after resistance and weakens assurance after compliance. The action replaced, not the operation alone, therefore determines how the target evaluates resistance and compliance.

Greater reach or lower relative cost makes the operation more attractive after either response. A larger audience cost also favors the operation after resistance because using it avoids the cost of backing down. A larger concession favors the operation over war after compliance because the operation preserves the transfer and war does not.

Having established when the coercer chooses war, the operation, or settlement, I turn to the target's response. The target compares the expected punishment from resistance with the expected cost of compliance. Compliance can cost it both the concession and any harm from force that follows. It accepts when that combined cost is no greater than the punishment from resistance. The largest accepted demand makes the target indifferent. [SI Lemma 4](#) states the full acceptance condition.

To obtain a closed-form solution, impose three additional conditions.

ASSUMPTION 2 (Closed-Form Region): At the selected demand x^* , (i) the war cutoffs and the post-compliance operation–settlement cutoff are interior; (ii) the operation is active after compliance; and (iii) every nonfighting type uses the operation after resistance, so $\kappa_{L,R} \geq \bar{c}$, equivalently $z + a \geq \lambda \bar{c}$.

Each condition has a separate role. Interior cutoffs make action probabilities linear under the uniform distribution. Activity after compliance ensures that all three continuation actions matter. The condition on $\kappa_{L,R}$ sets $\Pr^*(\text{settlement} \mid r = 1) = 0$, so every type that does not fight after resistance uses the operation. The interior post-compliance operation–settlement cutoff implies $\Pr^*(\text{settlement} \mid x^*) > 0$. A positive probability of settlement after compliance leaves some types that use the operation after resistance willing to settle after compliance. Compliance protects the target from the operation for exactly these

types. [Assumption 1\(i\)](#) keeps $S(y) = z$ at accepted demands. These restrictions produce the closed-form solution. The threat–assurance identity, the common threshold for payoff loss and war, and the expected-payoff decomposition hold beyond this region; complete coercive failure requires only [Assumption 1](#).

PROPOSITION 2 (Pooling Sustains the Largest Concession with Subwar Capability): Under [Assumptions 1](#) and [2](#), a pooling equilibrium exists in which every type demands x^* and war occurs with probability $\Pr^*(\text{war} \mid x^*)$. This equilibrium attains the largest demand sustainable under pure target responses and satisfies the intuitive criterion. At that demand, type payoffs and unconditional military outcomes are unique.

The formula for x^* is derived in [SI §C.3. Proof of Propositions 1 and 2](#). On this branch, some types use force after resistance but settle after compliance, giving the target a reason to concede. If $\Pr^*(\text{settlement} \mid x^*) = 0$, compliance provides no protection from the operation and the target accepts no positive demand.

Consider the operation–settlement choice after compliance. In the benchmark, a larger transfer raises the settlement payoff but leaves the war payoff unchanged. It therefore makes settlement more likely. With the operation, [Assumption 1\(i\)](#) gives $S(y) = z$ at every accepted demand. Settlement pays y , while the operation pays $y + z - \lambda c_1$, so their difference does not depend on the transfer. [SI §D.2. Settlement After Compliance](#) gives the formal statement and proof.

While the operation is active, increasing the concession can make some types switch from war to the operation. It does not change which types settle. After resistance, adding the operation can replace war, reducing the probability of war. But it can also replace settlement, so the target need not face less punishment overall. The target therefore weighs the expected cost of conceding, including any force that follows, against the expected punishment from resistance. Let H_R^* denote expected punishment after resistance and H_C^* expected harm from war or the operation after compliance with x^* .

REMARK 1 (Coercive Ability as Threat Minus Assurance Failure): If a common accepted demand leaves the target indifferent, coercive ability equals expected punishment after resistance, H_R^* , minus harm not prevented by compliance, H_C^* . It is zero exactly when $H_R^* = H_C^*$. This identity follows by rearranging the target's binding acceptance condition and requires neither the closed-form equilibrium nor a uniform cost distribution.

$$\begin{aligned} \Omega^* = & \underbrace{(p + c_2) \Pr^*(\text{war} \mid r = 1) + z \Pr^*(\text{operation} \mid r = 1)}_{H_R^*: \text{punishment after resistance}} \\ & - \underbrace{[(p + c_2) \Pr^*(\text{war} \mid x^*) + S(x^*) \Pr^*(\text{operation} \mid x^*)]}_{H_C^*: \text{harm not prevented by compliance}} \end{aligned} \quad (3)$$

Each component adds expected war harm and expected subwar harm, evaluated after resistance for H_R^* and after compliance for H_C^* . War causes loss $p + c_2$; the operation causes loss z after resistance and $S(x^*)$ after compliance.

Equation (3) is obtained by rearranging the target's binding acceptance condition: the expected concession it is willing to give up equals the punishment it avoids, less the harm it still expects after compliance. SI §D.1. Proof of Remark 1 gives the derivation.

Let $H_R^\dagger$ and $H_C^\dagger$ be the corresponding benchmark quantities in (2). The operation increases coercive ability when its effect on punishment exceeds its effect on harm after compliance, and reduces coercive ability when the reverse holds. It weakens the threat by lowering H_R^* and weakens assurance by raising H_C^* . Which effect occurs depends on whether the operation replaces war or settlement after each target response. Removing the operation from both continuation menus recovers (2). Setting its gain to zero alone need not reproduce the benchmark, because an operation can still avoid the audience cost after resistance.

WHEN THREAT AND ASSURANCE WEAKEN TOGETHER

To isolate how the operation changes the coercer's use of force, first hold transfer y fixed and compare harm excluding the retained concession, the component measured by H_R and H_C . Types that replace war with the operation change this harm from $p + c_2$ to z ; types that replace settlement with it add harm z . At resistance, where $y = 0$, this continuation-harm component equals the target's total loss.

With interior cutoffs and a uniform cost distribution, these two groups have a simple relation: for every type shifted out of war, $\frac{1-\lambda}{\lambda}$ types are shifted out of settlement. This continuation-harm component therefore rises when $z > \lambda(p + c_2)$, because the harm added by former settlers dominates, and falls when $z < \lambda(p + c_2)$, because replacing war dominates. The equality $z = \lambda(p + c_2)$ marks the balance between the two groups. SI §D.13. Interior Switch Accounting gives the cutoff algebra.

The same logic applies after resistance and compliance, but the number of types switching at each history differs. The audience cost makes the operation more attractive after resistance because it avoids the penalty for backing down. The accepted transfer makes it more attractive after compliance because the coercer keeps what the target has already paid. The net effect on coercive ability depends on how these different groups change punishment and post-compliance harm.

In the closed-form region, the operation replaces settlement for every type that would otherwise settle after resistance, while some types still settle after compliance. There are therefore no remaining settlers whose switch could add punishment after resistance. This limit breaks the common ratio between the two switching groups: the operation can reduce punishment after resistance even when it increases harm after compliance.

The joint-loss condition in Proposition 3 requires both effects. Its first inequality requires the harm saved by replacing war after resistance to exceed the harm added by replacing settlement. Its second requires the operation's direct effect and the concession response together to raise harm after compliance. When $z < \lambda(p + c_2)$, the operation

reduces harm at a fixed concession, so joint deterioration requires the concession reduction to overturn that benefit. When $z > \lambda(p + c_2)$, the operation directly raises harm after compliance. SI §D.13. [Interior Switch Accounting](#) gives the exact inequalities.

The link between the two components runs in one direction. Punishment after resistance is evaluated before any concession, so it does not vary with the demand. A weaker threat, however, lowers the largest accepted demand. The smaller concession makes war after compliance more attractive and raises the harm that compliance fails to prevent. This additional harm further reduces what the target will concede.

To isolate this feedback, hold the post-compliance continuation rules fixed as functions of the demand. In the closed-form region, the concession falls by more than the initial punishment loss when $x^* < p + c_2 - z$, meaning war is more costly for the target than conceding and then facing the operation. SI §D.1. [Proof of Remark 1](#) derives this response. Changing an operation parameter can also change the continuation rules, so this comparison isolates the effect of punishment.

Coercive Ability

When punishment falls and post-compliance harm rises, both changes reduce the concession that survives without war. More generally, coercive ability falls whenever the change in post-compliance harm exceeds the change in punishment, as (3) implies. Joint deterioration is therefore sufficient, although a large assurance loss can also outweigh a stronger threat.

PROPOSITION 3 (Subwar Capability Weakens Threats, Assurances, and Coercive Ability): Suppose the benchmark war cutoffs are interior at every accepted demand and the conditions for Proposition 2 hold. Threat and assurance both strictly weaken if and only if the two inequalities in SI [the joint-loss condition](#) hold. In that case, coercive ability is strictly lower. For every fixed set of the remaining primitives admitting these conditions, joint deterioration holds for all sufficiently

high audience costs. More generally, coercive ability falls exactly when the change in harm after compliance exceeds the change in punishment after resistance.

The conditions therefore identify a family of parameter intervals across military power, target war costs, and operation technologies. A higher cost of backing down strengthens the benchmark threat and raises its accepted concession. Once every nonfighting type operates after resistance, however, that cost no longer affects behavior in the operation game. Both components of coercion then deteriorate relative to the benchmark throughout an upper interval of this political cost. SI §D.10. [A Common Interval of Adverse Outcomes](#) proves that the payoff, war, and military-power results below hold together on such an interval.

Subwar capability also changes how military power translates into coercive ability, holding operation reach and relative cost fixed. When $z < \lambda\bar{c} \leq z + a$, no coercer type settles after resistance with the operation. The cost of backing down then ceases to affect the selected concession in that game. It continues to strengthen the benchmark's response to military power.

PROPOSITION 4 (Subwar Capability Reduces the Coercive Return to Military Power): Maintain Assumption 1, interior benchmark war cutoffs, and $z < \lambda\bar{c} \leq z + a$. At values of p where derivatives exist, military power weakly increases Ω^* and strictly increases $\Omega^\dagger$. Holding the other primitives fixed, a unique algebraic threshold $a_P(p) \geq 0$ characterizes suppression among the admissible values of a :

$$0 \leq \frac{\partial \Omega^*}{\partial p} < \frac{\partial \Omega^\dagger}{\partial p} \quad \text{if and only if} \quad a > a_P(p).$$

The difference between the benchmark and operation-game slopes strictly increases with a . If the active, interior conditions of Proposition 2 also hold, suppression occurs for all sufficiently high audience costs.

The threshold is defined in primitives in SI §D.9. [Proof of Proposition 4](#). On the active, interior branch, it lies strictly below the upper admissible cost of backing down, guaran-

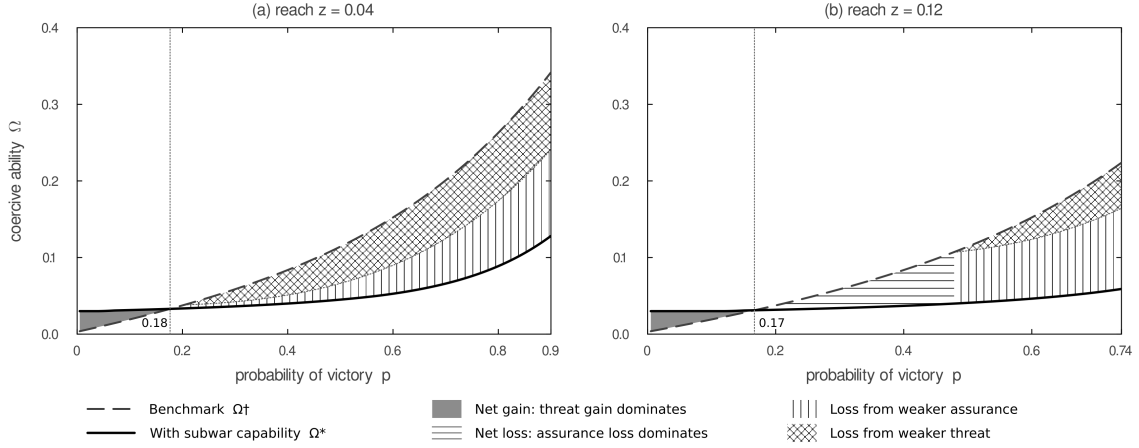


Figure 2: Subwar Capabilities Weaken the Effect of Military Power on Coercive Ability

Note: Coercive ability in the benchmark (dashed) and with the subwar capability (solid) as the coerced's probability of victory varies. The shaded gap is the difference between the two curves. Grey shading marks a net gain when the gain in threat outweighs the loss of assurance; horizontal lines mark a net loss when the assurance loss outweighs the threat gain. Where both weaken, vertical lines and crossed diagonal lines divide the loss into weaker assurance and weaker threat, respectively. Parameters: $c_2 = 0.02$, $a = 0.30$, $\lambda = 0.08$, c_1 uniform on $[0, 2]$; $z = 0.04$ in panel (a) and $z = 0.12$ in panel (b).

teeing an interval of suppression for every feasible configuration of the other parameters. Proposition 4 also covers other action regimes, where the threshold may lie outside the admissible range.

Both panels of Figure 2 satisfy the slope ordering in Proposition 4 throughout their plotted ranges, including one-sided derivatives where the coerced's action changes. SI §D.9. Proof of Proposition 4 verifies this comparison analytically. The subwar capability can therefore raise coercive ability at low military power while reducing the gain from additional power. The curves cross while punishment after resistance is still greater with the capability: assurance failure first outweighs the stronger threat, and weaker punishment later compounds the loss.

The two panels also differ in what produces the loss. At the smaller reach, no coerced type eventually chooses the operation after compliance; beyond that point the loss comes from weaker punishment after resistance together with the war that the smaller demand

no longer prevents. At the larger reach, post-compliance harm combines the operation's direct effect with additional war from the smaller demand. Proposition 3 establishes joint deterioration over intervals of political costs; the panels trace how its two components change as military power varies.

The exact comparison, including cases where threat and assurance move in opposite directions, is derived in SI §D.8. [Proof of Proposition 3](#). Its mechanism can be seen by separating the protection that compliance buys in the closed-form region:

$$\begin{aligned} \Omega^* = & \underbrace{(p + c_2 - z)(\Pr^*(\text{war} \mid r = 1) - \Pr^*(\text{war} \mid x^*))}_{\text{protection from war}} \\ & + \underbrace{z \Pr^*(\text{settlement} \mid x^*)}_{\text{protection from the operation}} \end{aligned} \tag{4}$$

The first term values the reduction in war probability by the harm war adds beyond the operation. The second values the chance of settlement, which avoids the operation's harm as well. These two forms of protection determine how much the target will concede.

Greater reach makes avoiding the operation more valuable, but it also makes the coercer less willing to settle after compliance. The second term can therefore first rise and then fall as reach grows. Because reach changes both forms of protection, its overall effect on coercive ability cannot be inferred from the operation's harm or the war probability alone.

In the closed-form region, protection from the operation, $z \Pr^*(\text{settlement} \mid x^*)$, is independent of military power. The effect of power on total coercive ability also accounts for changes in the accepted concession and protection from war.

The operation–settlement cutoff identifies complete coercive failure.

REMARK 2 (Complete Coercive Failure): Under Assumption 1, coercive ability in the selected pooling equilibrium is zero if and only if $z \geq \lambda \bar{c}$. When the inequality holds, the operation–settlement cutoff after compliance is at or above the top of the type distribution, so even the highest-cost coercer type weakly prefers the operation to settlement after compliance and chooses it at a tie. Compliance then provides no protection from the operation, $\Pr^*(\text{settlement} \mid x^*) = 0$, the target refuses every positive demand, and $x^* = \Omega^* = 0$.

Complete failure occupies the entire region $z \geq \lambda \bar{c}$, for every positive cost of backing down, and requires no closed-form or interior-cutoff restrictions. For any positive operational reach, sufficiently low operating cost therefore eliminates every positive concession while preserving Assumption 1. Moreover, every equilibrium with pure target responses yields zero retained concession in this region, as SI Lemma 5 shows. Below the threshold, a positive measure of types uses force after resistance but settles after compliance with zero, making a sufficiently small concession acceptable. At or above it, neither response to zero produces settlement, and every positive demand is rejected. The active-branch formula $\Pr^*(\text{settlement} \mid x^*) = 1 - \frac{z}{\lambda \bar{c}}$ is therefore replaced by zero when its cutoff reaches the top of the support.

The Coercer’s Payoff

The concession response connects the coercer’s private loss to the prospect of war. The same threshold determines when every coercer type is weakly worse off and when war is weakly more likely, even outside the closed-form region. It compares the concession lost with the operation’s net gain for the benchmark type indifferent between war and settlement. SI Proposition 9 gives the formal condition.

The benchmark type at the war margin has the most to gain from the operation relative to its benchmark payoff. Lower-cost types already prefer war; higher-cost types pay more to operate. Once the lost concession exhausts that marginal type’s net gain, no type benefits

from the expanded menu. The same comparison determines whether the operation still prevents war among benchmark fighters. At the threshold, war is unchanged; above it, some benchmark settlers fight. The equivalence applies whether the operation remains active after compliance or is dominated, without the closed-form restrictions. At the two accepted demands, the payoff, war, and conflict-cost implications also hold for any continuous cost distribution with full support. When concessions shift military power or remain partly secure during war, the payoff-loss and war thresholds generally differ; SI §F.2. [Power Shift and Partial Recontestation](#) gives the separate conditions.

Figure 3 shows why accounting for the operation’s cost matters. At the illustrated $p = 0.55$, the lost concession is smaller than the operation’s gross gain but exceeds its net gain for the benchmark type at the war margin. Every benchmark settler therefore loses, and war rises. Types that fight in both games are unaffected. The payoff loss can come from switching from settlement to war, paying to operate, or settling for a smaller concession.

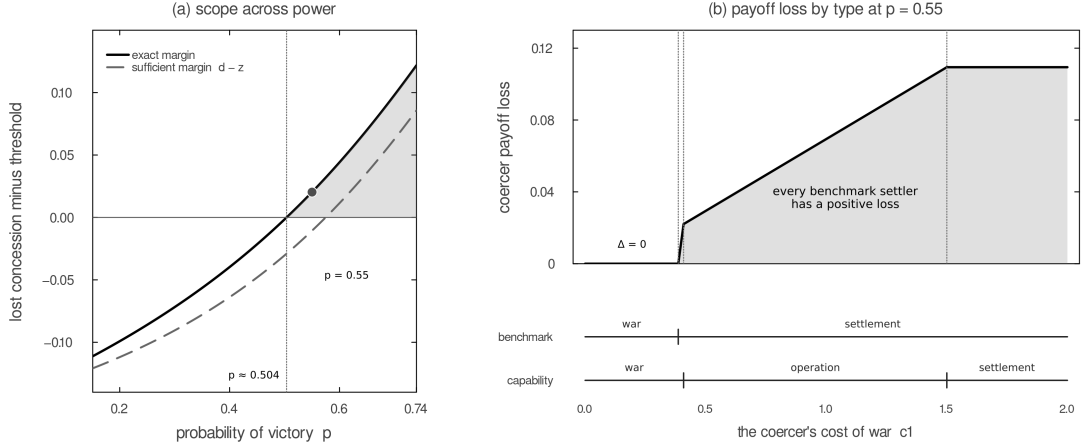


Figure 3: The Common Threshold for Payoff Loss and War

Note: Panel (a) plots the exact margin $x^\dagger - x^* - \max\{0, z - \lambda(p - x^\dagger)\}$ as a solid curve and the stronger sufficient margin $x^\dagger - x^* - z$ as a dashed guide. The exact margin crosses zero at $p \approx 0.504$, while the sufficient margin crosses at $p \approx 0.575$. Panel (b) fixes $p = 0.55$ and plots the benchmark payoff minus the operation-game payoff by coerced type; the aligned intervals show each game’s post-compliance action. Parameters: $c_2 = 0.02$, $a = 0.30$, $\lambda = 0.08$, $z = 0.12$, and c_1 uniform on $[0, 2]$.

These losses and the increase in war occur together throughout the common interval established above. For every feasible active, interior operation profile, a sufficiently high audience cost makes the lost concession exceed even the operation's gross gain. Every benchmark settler then strictly loses, the coercer's expected payoff falls, and both war and joint conflict cost rise.

Proposition 5 separates the two effects on the coercer's expected payoff: gaining another military option at a fixed concession and receiving a different concession from the target.

PROPOSITION 5 (Concession Loss Outweighs the Gain from Subwar Capability):

The change in expected payoff equals the gain from adding the operation at the new concession plus the payoff effect of changing the concession in the benchmark. The first component is nonnegative. The second is negative when $x^* < x^\dagger$. Expected payoff falls exactly when the payoff loss from the smaller concession outweighs the gain from the additional option.

Holding the concession fixed, the operation cannot hurt the coercer because every original action remains available. A smaller accepted concession can nevertheless make possession of the capability costly overall. This decomposition applies across all action regimes and does not require uniformity. Only the displayed closed-form payoffs in SI §D.4. [Proof of Proposition 5](#) require uniformity, interior cutoffs, and an active operation.

SUBWAR CAPABILITY AND THE PROSPECT OF PEACE

The operation changes both the available continuation actions and the concession at which the coercer chooses among them. At a fixed concession, it can replace war among low-cost types or settlement among intermediate types below κ_L . When the accepted demand falls, however, some benchmark settlers choose war. These types never use the operation. [Figure 4](#) compares the resulting post-compliance actions across the two games.

The same threshold applies whether the operation remains active after compliance or becomes dominated. War must be preferable to both alternatives. While the operation is

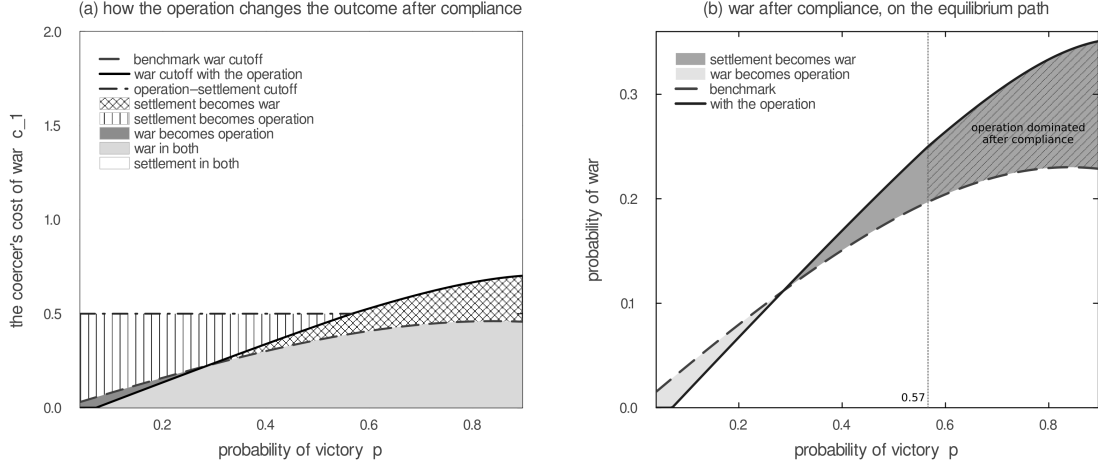


Figure 4: Subwar Capability Can Replace War Yet Increase It

Note: Panel (a) uses gray fills and hatching for regions bounded by the analytical equilibrium cutoffs, using piecewise closed-form demands at the calibration of Figure 2 with $z = 0.04$. Dashed and solid lines show the benchmark and operation-game war cutoffs; the dash-dotted line marks the operation-settlement cutoff while the operation is active. The operation region's upper boundary is horizontal because $\kappa_L = \frac{z}{\lambda}$ contains neither p nor the demand. Panel (b) plots the probability of war after compliance on each game's equilibrium path: below $p \approx 0.29$ the operation replaces war; above it, the smaller accepted demand pushes types that would have kept the benchmark settlement into war. In panel (b), diagonal hatching marks the range above $p \approx 0.57$, where the operation is dominated after compliance. There the post-compliance menu is the benchmark menu evaluated at a smaller demand, and the entire increase follows from the weaker punishment the target expects after resistance.

active, it sets the relevant war cutoff; otherwise settlement does. In either case, war rises when the smaller concession makes some benchmark settlers prefer to fight.

Panel (b) distinguishes two ranges in which war exceeds its benchmark probability. Between the crossing at $p \approx 0.29$ and the dotted line at $p \approx 0.57$, the operation remains active after compliance. Its anticipated use weakens assurance directly. In the hatched range above $p \approx 0.57$, the operation is dominated after compliance and is never used on the equilibrium path. It nevertheless replaces war after resistance, weakening punishment and reducing the concession. In both ranges, the smaller concession draws benchmark settlers into war and thereby weakens assurance further. These effects are jointly determined in equilibrium, not a temporal sequence. Additional war consists of direct choices of war by types that never use the operation.

Under Assumption 1, suppose some coercer types would use the operation if the target resisted, but every type chooses war or settlement after the selected concession. The operation is then never used in equilibrium, yet its availability weakens threats and assurances, lowers the concession and coercive ability, leaves every coercer type weakly worse off, and strictly increases war and conflict cost. SI §D.12. [An Unused Operation](#) proves these conclusions without additional interior-cutoff or backing-down-cost restrictions. Possession can therefore affect bargaining through the punishment expected after resistance even when no operation is observed after compliance. An analysis based only on observed operations would miss that channel.

Subwar Capability and the Cost of Conflict

More efficient limited operations would seem to promise less costly conflict. Yet making subwar operations less costly to conduct can increase the expected joint cost of conflict by weakening a government's coercive ability. Let K denote expected joint conflict cost in the model. Settlement costs nothing, war costs the two states $c_1 + c_2$, and the operation costs λc_1 . The baseline treats both the demand and the operation's gain as transfers between the states, so modeled joint welfare is $1 - K$. This accounting excludes any physical destruction whose loss to the target is not offset by the coercer's modeled gain.

Concession loss. As the accepted concession falls, expected conflict cost can exceed its benchmark level even while war remains less likely than in the benchmark. To isolate this ordering, hold the operation's primitives fixed and vary the concession lost. Let $b = p - x^\dagger$ be the benchmark war cutoff and $\Delta K(d)$ operation-game conflict cost after concession $x^\dagger - d$ minus benchmark conflict cost. Thus $\Delta K(x^\dagger - x^*) = K^* - K^\dagger$, while $\Delta K(0)$ is the direct cost effect at the benchmark concession. Along $0 \leq d \leq x^\dagger$, require an interior benchmark war cutoff and interior relevant continuation cutoffs at every evaluated transfer, including $d = 0$. The operation may leave the menu, provided the remaining war-settlement cutoff stays interior. Suppose the operation reaches some benchmark settlers, $\frac{z}{\lambda} > b$, and define the war threshold $d_W = z - \lambda b$ with $0 < d_W < x^\dagger$.

Under these path conditions, Proposition 6 establishes that conflict cost strictly increases with the lost concession. The ordering holds for any cost distribution with positive density over the relevant type intervals. While the operation is active, a smaller concession moves types from the operation into war; once it is dominated, a further reduction moves types from settlement into war. If $\Delta K(0) < 0$, the initial saving is exhausted at a unique $d_K \in (0, d_W)$. Cost rises above its benchmark level before war does because, when war returns to its benchmark probability, the operation still replaces settlement for some types and adds λc_1 in costs.

Lower operating cost. Reducing the operation's cost changes both its use and the concession. These comparisons are local and do not require the whole lost-concession path to have interior cutoffs. Hold p, c_2, a, z , and $\bar{c}$ fixed on an active interior post-compliance branch. At a fixed concession, lower operating cost saves resources on operations already conducted and encourages additional use. Operations replacing war can save costs; operations replacing settlement add costs. Total conflict cost rises when those added costs outweigh the savings from reduced operating cost and avoided war. SI §D.5. Proof of Proposition 6 gives the exact condition.

When the concession adjusts, maintain Proposition 2's closed-form conditions with $z + a > \lambda \bar{c}$, so no type settles after resistance and this restriction is strict.

PROPOSITION 6 (Lower Operation Costs Increase Total Conflict Cost): *Concession loss.* Under the stated path conditions, $\Delta K(d)$ strictly increases and, if $\Delta K(0) < 0$, crosses zero at a unique $d_K \in (0, d_W)$.

Lower operating cost. In the stated closed-form region, a reduction in operating cost lowers the accepted demand. It raises equilibrium conflict cost whenever it already raises cost at the fixed equilibrium concession. This sufficient condition always holds when $c_2 = 0$. More generally, if Assumption 1(i) holds and $\bar{c} - z > 2p + c_2$, both effects hold on a nonempty interval immediately above $\lambda = \frac{z}{\bar{c}}$, including when $c_2 > 0$.

On the positive-war branch, reducing operating cost shifts types from war to the operation after resistance and reduces punishment. After compliance, it can shift types from war or settlement to the operation, with opposing effects on protection. Accounting for both substitutions, the positive demand derivative establishes that the target concedes less as λ falls. That response moves additional types toward war. Throughout [Figure 5](#)'s range, the settlement substitution also dominates the direct post-compliance comparison, so both punishment and assurance weaken.

The smaller equilibrium concession raises cost further. Whenever the added operations outweigh the direct cost savings at a fixed concession, the bargaining response compounds the increase. The last part of [Proposition 6](#) establishes an interval of rising conflict costs as operating cost approaches the point of complete coercive failure, for any target war cost satisfying the stated bounds. On the positive-war branch, the concession response is strong enough near this threshold even when replacing war saves the target substantial costs. If war has already disappeared, added operations replacing settlement drive the increase instead. As λ falls from 0.21 to 0.061 in [Figure 5](#), cost at the benchmark concession rises from 0.051 to 0.087. Allowing the concession to adjust raises equilibrium cost from 0.057 to 0.138, while benchmark cost remains 0.050.

SCOPE OF THE EQUILIBRIUM COMPARISON

Part of the analysis does not depend on equilibrium selection. If the operation takes a fixed amount and is active after an accepted demand, the coercer settles if and only if $c_1 > \frac{z}{\lambda}$. That cutoff does not depend on the demand or off-path beliefs, and [SI Proposition 8](#) proves it for every equilibrium. The threat-minus-assurance identity in [Remark 1](#) requires only a common accepted demand that leaves the target indifferent. The cutoff result additionally requires the operation to be active at that demand. Different equilibria can place different demands on the path and associate different types with them, so the cutoff alone does not determine the unconditional settlement rate or coercive ability.

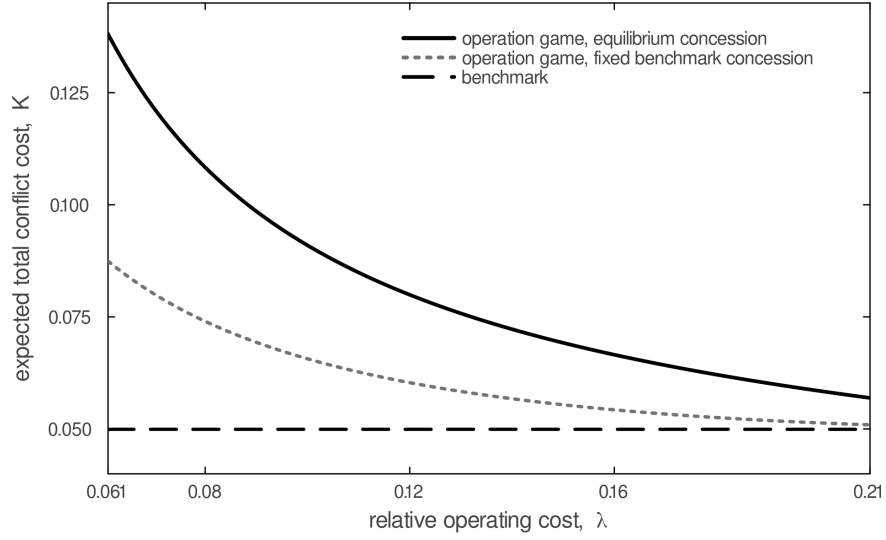


Figure 5: Lower Operating Cost Can Raise Total Conflict Cost

Note: The solid curve plots equilibrium expected total conflict cost with the operation. The dotted gray curve holds the concession at its benchmark level, isolating the direct effect of changing λ . The dashed black line is benchmark cost. Lower λ means a lower cost of conducting the operation relative to war. Over the displayed interval, lowering λ weakens the threat, weakens assurance on net, and reduces the accepted concession. The interval satisfies the closed-form premises and the fixed-concession cost condition in SI §D.5. [Proof of Proposition 6](#) Proposition 6 gives the positive demand derivative. The upper endpoint $\lambda = 0.21$ satisfies $z + a = \lambda \bar{c}$, so its derivative comparison is from below. Parameters: $p = 0.65$, $c_2 = 0.02$, $a = 0.30$, $z = 0.12$, and c_1 uniform on $[0, 2]$.

Pure target responses admit an interval of largest accepted demands. SI [Lemma 5](#) characterizes this interval under Assumption 1 and shows that every demand in it supports a pooling equilibrium satisfying the intuitive criterion. Types that strictly benefit from compliance choose the largest accepted demand; indifferent types may send other signals. At a fixed largest demand, those signals preserve type payoffs and unconditional military outcomes.^[6]

The selected equilibrium maximizes every coercer type's payoff and the unconditional retained concession. It also minimizes unconditional war and expected joint conflict cost

^[6]They can change the probability of compliance and the posterior at an accepted demand. War and retained concessions conditional on compliance can therefore differ across signal strategies. The reported coercive-ability formulas describe the selected pooling assessments.

within the pure-response class, as SI [Corollary 2](#) shows. A larger concession raises the payoffs from settlement and the operation relative to war. Thus, when the capability raises war or conflict cost in the selected comparison, it raises the corresponding minimum across equilibria. The comparison does not rank every pair of equilibria across the games or establish which equilibrium an actual crisis selects.

Zero demand can coexist with positive accepted demands. In the benchmark, it is an equilibrium exactly when $p + a < \bar{c}$. In the operation game’s region without settlement after resistance, it is an equilibrium exactly when $z > \lambda p$. If $z \geq \lambda \bar{c}$, every pure-response equilibrium instead yields zero retained concession.^[7]

CONCLUSION

A subwar capability can undermine the bargain that makes restraint worthwhile. A subwar operation can weaken punishment after resistance and protection after compliance, reducing the target’s concession. The smaller concession makes war more attractive after compliance. These adverse effects hold together over intervals of political costs across military and technological configurations. Within the baseline payoff structure, the same threshold determines when every coercer type is weakly worse off and when war is weakly more likely, without requiring the closed-form equilibrium or a uniform cost distribution. Strict coercive losses and increases in war and conflict cost also survive nearby nonuniform priors and sufficiently small escalation risk under the conditions established in the Supporting Information.

These results change what observed force can reveal about capability. An operation that goes unused can still weaken the threat after resistance, reduce the concession, and induce war among benchmark settlers. Conversely, fewer general wars need not imply lower modeled conflict costs: when an initial operating-cost saving exists under the concession-path conditions, it disappears before war returns to its benchmark rate. Studying capability

^[7]SI [§E.1. Zero-Demand Equilibria](#) gives the full conditions, including settlement after resistance and cases in which zero is the only attainable demand.

therefore requires attention to both the bargains it changes and the forms of force that follow them.

These results suggest why a subwar capability’s value in deterrence need not carry over to compellence. Both require punishment to be conditional, but the compellence problem here asks the target to make a concession and then rely on the coercer not to use the operation. Because the operation remains available after compliance, it can strengthen punishment after resistance while weakening the protection that compliance buys. A subwar option may therefore support a deterrent threat yet weaken a compellent bargain. The effect of a subwar capability on compellence cannot be inferred from threat credibility alone (Schelling 1966; Kydd and McManus 2017; Cebul, Dafoe, and Monteiro 2021).

The analysis also qualifies the usual option-value logic. At a fixed transfer, adding the operation cannot hurt the coercer because every benchmark action remains available. When the coercer makes the largest demand the target will accept, however, the expanded menu can reduce the transfer and make the coercer worse off. This suggests a role for commitments that preserve punishment after resistance while limiting the operation after compliance. Such a restriction could improve the earlier bargain by protecting the target after it concedes, an application of the power to bind oneself (Schelling 1960).^[8]

^[8]I used OpenAI Codex (recorded model identifiers: gpt-5.6-sol and gpt-6-astra; July–September 2026) and Anthropic Claude (claude-opus-5, claude-fable-5, and claude-fable-5-1; August–September 2026) for editorial support and proof drafting. Editorial assistance included copyediting, formatting, citation checks, and mathematical notation. AI contributions were thus limited to editorial support and proof drafting. I personally checked all AI-drafted proofs and reviewed the retained editorial changes. I take responsibility for the manuscript’s arguments, mathematical results, citations, and factual accuracy.

REFERENCES

- Altman, Dan. 2017. "By Fait Accompli, Not Coercion: How States Wrest Territory from Their Adversaries." *International Studies Quarterly* 61(4): 881–91. doi:[10.1093/isq/sqx049](https://doi.org/10.1093/isq/sqx049).
- Altman, Dan. 2020. "The Evolution of Territorial Conquest After 1945 and the Limits of the Territorial Integrity Norm." *International Organization* 74(3): 490–522. doi:[10.1017/s0020818320000119](https://doi.org/10.1017/s0020818320000119).
- Art, Robert J., and Kelly M. Greenhill. 2018. "The Power and Limits of Compellence: A Research Note." *Political Science Quarterly* 133(1): 77–97. doi:[10.1002/polq.12738](https://doi.org/10.1002/polq.12738).
- Banks, Jeffrey S. 1990. "Equilibrium Behavior in Crisis Bargaining Games." *American Journal of Political Science* 34(3): 599. doi:[10.2307/2111390](https://doi.org/10.2307/2111390).
- Cebul, Matthew D., Allan Dafoe, and Nuno P. Monteiro. 2021. "Coercion and the Credibility of Assurances." *The Journal of Politics* 83(3): 975–91. doi:[10.1086/711132](https://doi.org/10.1086/711132).
- Cho, In-Koo, and David M. Kreps. 1987. "Signaling Games and Stable Equilibria." *The Quarterly Journal of Economics* 102(2): 179–221. doi:[10.2307/1885060](https://doi.org/10.2307/1885060).
- Fearon, James D. 1994. "Domestic Political Audiences and the Escalation of International Disputes." *American Political Science Review* 88(3): 577–92. doi:[10.2307/2944796](https://doi.org/10.2307/2944796).
- Fearon, James D. 1997. "Signaling Foreign Policy Interests." *Journal of Conflict Resolution* 41(1): 68–90. doi:[10.1177/0022002797041001004](https://doi.org/10.1177/0022002797041001004).
- Gannon, J. Andrés, Erik Gartzke, Jon R. Lindsay, and Peter Schram. 2024. "The Shadow of Deterrence: Why Capable Actors Engage in Contests Short of War." *Journal of Conflict Resolution* 68(2–3): 230–68. doi:[10.1177/00220027231166345](https://doi.org/10.1177/00220027231166345).
- Gat, Azar. 2013. "Is War Declining – and Why?." *Journal of Peace Research* 50(2): 149–57. doi:[10.1177/0022343312461023](https://doi.org/10.1177/0022343312461023).

- Gleditsch, Kristian Skrede, and Steve Pickering. 2013. “Wars Are Becoming Less Frequent: A Response to Harrison and Wolf.” *The Economic History Review* 67(1): 214–30. doi:10.1111/1468-0289.12002.
- Horowitz, Michael C. 2020. “Do Emerging Military Technologies Matter for International Politics?.” *Annual Review of Political Science* 23(1): 385–400. doi:10.1146/annurev-polisci-050718-032725.
- Ioku, Shusuke, and Shuhei Kurizaki. 2026. “Assurance in Coercive Bargaining.” Working paper. https://shusuke-ioku.github.io/files/attack_option_paper.pdf.
- Jensen, Benjamin M, Christopher Whyte, and Scott Cuomo. 2019. “Algorithms at War: The Promise, Peril, And Limits of Artificial Intelligence.” *International Studies Review* 22(3): 526–50. doi:10.1093/isr/viz025.
- Kenkel, Brenton, and Peter Schram. 2025. “Uncertainty in Crisis Bargaining with Multiple Policy Options.” *American Journal of Political Science* 69(1): 194–209. doi:10.1111/ajps.12849.
- Kydd, Andrew H., and Roseanne W. McManus. 2017. “Threats and Assurances in Crisis Bargaining.” *Journal of Conflict Resolution* 61(2): 325–48. doi:10.1177/0022002715576571.
- Maass, Richard W. 2021. “Salami Tactics: Faits Accomplis and International Expansion in the Shadow of Major War.” *Texas National Security Review* 5(1): 33–54. <https://tnsr.org/2021/11/salami-tactics-faits-accomplis-and-international-expansion-in-the-shadow-of-major-war/>.
- Paine, Jack, and Scott A. Tyson. 2020. “Uses and Abuses of Formal Models in Political Science” eds. Dirk Berg-Schlosser, Bertrand Badie, and Leonardo Morlino. *The SAGE Handbook of Political Science: A Global Perspective*: 188–202. doi:10.4135/9781529714333.n14.

- Pape, Robert A. 1996. *Bombing to Win: Air Power and Coercion in War*. Ithaca, NY: Cornell University Press. doi:[10.7591/9780801471513](https://doi.org/10.7591/9780801471513).
- Pauly, Reid B. C. 2024. “Damned If They Do, Damned If They Don't: The Assurance Dilemma in International Coercion.” *International Security* 49(1): 91–132. doi:[10.1162/isec_a_00488](https://doi.org/10.1162/isec_a_00488).
- Pauly, Reid B. C. 2025. *The Art of Coercion*. Cornell University Press. doi:[10.1353/book.131106](https://doi.org/10.1353/book.131106).
- Pinker, Steven. 2011. *The Better Angels of Our Nature: The Decline of Violence in History and Its Causes*. London: Allen Lane.
- Powell, Robert. 1989. “Nuclear Deterrence and the Strategy of Limited Retaliation.” *American Political Science Review* 83(2): 503–19. doi:[10.2307/1962402](https://doi.org/10.2307/1962402).
- Schelling, Thomas C. 1960. *The Strategy of Conflict*. Cambridge, MA: Harvard University Press.
- Schelling, Thomas C. 1966. *Arms and Influence*. New Haven, CT: Yale University Press. doi:[10.12987/9780300253481](https://doi.org/10.12987/9780300253481).
- Schram, Peter. 2021. “Hassling: How States Prevent a Preventive War.” *American Journal of Political Science* 65(2): 294–308. doi:[10.1111/ajps.12538](https://doi.org/10.1111/ajps.12538).
- Schram, Peter. 2022. “When Capabilities Backfire: How Improved Hassling Capabilities Produce Worse Outcomes.” *The Journal of Politics* 84(4): 2246–60. doi:[10.1086/719637](https://doi.org/10.1086/719637).
- Schultz, Kenneth A. 1999. “Do Democratic Institutions Constrain or Inform? Contrasting Two Institutional Perspectives on Democracy and War.” *International Organization* 53(2): 233–66. doi:[10.1162/002081899550878](https://doi.org/10.1162/002081899550878).
- Sechser, Todd S. 2010. “Goliath's Curse: Coercive Threats and Asymmetric Power.” *International Organization* 64(4): 627–60. doi:[10.1017/s0020818310000214](https://doi.org/10.1017/s0020818310000214).

- Slantchev, Branislav L. 2003. "The Power to Hurt: Costly Conflict with Completely Informed States." *American Political Science Review* 97(1): 123–33. doi:[10.1017/s000305540300056x](https://doi.org/10.1017/s000305540300056x).
- Slantchev, Branislav L. 2005. "Military Coercion in Interstate Crises." *American Political Science Review* 99(4): 533–47. doi:[10.1017/s0003055405051865](https://doi.org/10.1017/s0003055405051865).
- Slantchev, Branislav L. 2011. *Military Threats: The Costs of Coercion and the Price of Peace*. Cambridge: Cambridge University Press. doi:[10.1017/CBO9780511778940](https://doi.org/10.1017/CBO9780511778940).
- Spaniel, William. 2023. "Issue Indivisibility as a Cause of Peace." *The Journal of Politics* 85(2): 760–63. doi:[10.1086/723025](https://doi.org/10.1086/723025).
- Spaniel, William, and Işıl İdrisoğlu. 2024. "Endogenous Military Strategy and Crisis Bargaining." *Conflict Management and Peace Science* 41(4): 365–91. doi:[10.1177/07388942231180499](https://doi.org/10.1177/07388942231180499).
- Tarar, Ahmer. 2016. "A Strategic Logic of the Military*<i>fait Accompli*
- </i>." *International Studies Quarterly* 60(4): 742–52. doi:[10.1093/isq/sqw018](https://doi.org/10.1093/isq/sqw018).

Zegart, Amy. 2018. "Cheap Fights, Credible Threats: The Future of Armed Drones and Coercion." *Journal of Strategic Studies* 43(1): 6–46. doi:[10.1080/01402390.2018.1439747](https://doi.org/10.1080/01402390.2018.1439747).*

SUPPORTING INFORMATION

CONTENTS

A. Notation 2

B. Proofs of the Lemmas 2

C. Proof of the Equilibrium Characterization 5

D. Proofs of the Comparative Statics 6

E. Pure Target Responses and D1 16

F. Extensions and Scope 20

A. NOTATION

Write $C = \bar{c}$, $q = 1 - \lambda$, $h = p + c_2$, $X = x^\dagger$, and $Y = x^*$. A dagger denotes the benchmark; a star denotes the operation game. Histories are $(y, r) = (x, 0)$ after compliance and $(0, 1)$ after resistance. The operation gains $S(y) = \min\{z, 1 - y\}$ and costs λc_1 ; Assumption 1(i) gives $S(x) = z$ at accepted demands. Settlement pays $y - ar$, and war pays $p - c_1$ to the coercer and costs the target h relative to retaining the stake.

Under the prior, W is the probability of war and L of either military action: operation and settlement have probabilities $L - W$ and $1 - L$. Subscript R denotes resistance; superscript C denotes compliance (as in $W^C(0)$); a positive argument denotes compliance with that demand. Active-operation cutoffs give $W = F(\kappa_W)$ and $L = F(\kappa_L)$; otherwise $W = L = F(\kappa_S)$. Extend F by zero below 0 and one above C . At informative demands use the posterior instead. Write $\alpha^* = 1 - L(Y)$, $b = p - X$, and $d = X - Y$. Expected punishment and post-compliance harm are H_R, H_C ; expected coercer payoff and joint conflict cost are V, K . In the gain-harm extension, g is the coercer's gain and η the target's loss.

B. PROOFS OF THE LEMMAS

LEMMA 1 (The Coercer's Menu in the Benchmark): At (y, r) , war is optimal exactly when $c_1 \leq \kappa^\dagger(y, r) = p - y + a\mathbb{1}[r = 1]$; otherwise the coercer settles.

LEMMA 2 (No Type Prefers Resistance): For $x > 0$, concession weakly raises every type's payoff and strictly raises every nonfighter's payoff. Indifferent types are $[0, \kappa_W(x)]$ when the operation is active, with positive measure exactly when $p > x + S(x)$, and $[0, \kappa_S(x)]$ otherwise.

LEMMA 3 (The Coercer's Menu with a Subwar Operation): If $\kappa_W < \kappa_L$, war, the operation, and settlement are optimal in ascending cost intervals. Otherwise the operation is dominated.

$$\kappa_W(y) = \frac{p - y - S(y)}{1 - \lambda}, \quad \kappa_L(y, r) = \frac{S(y) + a\mathbb{1}[r = 1]}{\lambda} \quad (5)$$

Write $W = F(\kappa_W)$ and $L = F(\kappa_L)$ when active; otherwise $W = L = F(\kappa_S)$, where $\kappa_S = p - y + a\mathbb{1}[r = 1]$. Extend F by zero and one outside its support.

LEMMA 4 (The Target's Rule): The target concedes exactly when (6) holds, using the relevant posterior for the probabilities.

$$(p + c_2)W_R + z(L_R - W_R) \geq x(1 - W(x)) + (p + c_2)W(x) + S(x)(L(x) - W(x)) \quad (6)$$

B.1. Proof of Lemma 1

Proof. Comparing $p - c_1$ with $y - a\mathbb{1}[r = 1]$ gives the cutoff. $\square$

B.2. Proof of Lemma 2

Proof. War pays $p - c_1$ after either response. Settlement rises from $-a$ to x , and the operation from $z - \lambda c_1$ to $\min\{x + z, 1\} - \lambda c_1$, strictly higher since $x > 0$ and $z < 1$. Maximizing these payoffs proves weak dominance; equality requires war after both responses. The maintained continuation tie rule makes every nonfighter's gain strict. $\square$

B.3. Proof of Lemma 3

Proof. The three affine payoffs have slopes $-1, -\lambda, 0$. Pairwise indifference gives (5) and κ_S . The operation exceeds both alternatives exactly between κ_W and κ_L ; supported users have positive measure exactly when $\max\{0, \kappa_W\} < \min\{\bar{c}, \kappa_L\}$. For uncapped transfers, activity requires $z + a > \lambda(p + a)$ after resistance and $z > \lambda(p - x)$ after compliance. $\square$

B.4. Proof of Lemma 4

Proof. Expected target losses after resistance and compliance are respectively the left and right sides of (6). Comparing them proves the rule. $\square$

REMARK 3 (Distributional Scope): On the uncapped domain, the upper-tail hazard bound

$$\sup_{0 \leq c \leq \frac{p}{1-\lambda}} \frac{f(c)}{1 - F(c)} < \frac{1 - \lambda}{p + c_2}$$

makes the prior acceptance margin strictly decreasing. Under uniformity this is Assumption 1(ii). Cutoffs, the target rule, Propositions 7–8, and the threat–assurance identity require no uniformity; the quadratic demands do. Without this bound, connected acceptance sets are not asserted.

Proof. For $G(x) = V_2(x) - V_{2,R}$, differentiation gives $G' = f(p - x)(p + c_2 - x) - [1 - F(p - x)]$ on the dominated branch, and $G' = f(\kappa_W) \frac{p + c_2 - x - z}{1 - \lambda} - [1 - F(\kappa_W)]$ on the active branch. Nonpositive harm terms imply negativity directly; otherwise division by survival probability and the hazard bound does. Beyond the war region $G' = -1$, and branches join continuously. Uniformity gives supremum $\frac{1}{[\bar{c} - \frac{p}{1-\lambda}]}$, yielding Assumption 1(ii). $\square$

B.5. Robustness to Nearby Priors

Every strict comparison of demands, coercive ability, war, expected payoff, or joint conflict cost persists under sufficiently small smooth changes from the uniform prior, with common support and baseline demands strictly inside their continuation regimes.

Proof. Uniform closeness of CDFs and densities preserves the strict hazard bound because baseline survival is bounded away from zero on the relevant interval. Acceptance margins converge uniformly; their unique roots and strict regime inequalities persist. The belief construction below still applies. Probabilities are CDF evaluations at continuous cutoffs; payoffs and costs are integrals of bounded functions whose moving boundaries have zero probability. All therefore vary continuously, preserving strict signs, including Proposition 9’s strict threshold. Equality cases need not persist. $\square$

C. PROOF OF THE EQUILIBRIUM CHARACTERIZATION

C.1. Off-Path Beliefs and the Intuitive Criterion

Let $U_C(x, c_1)$ and $U_R(c_1)$ be optimized coerced payoffs. If U_C increases weakly with x , x^e is the largest prior-accepted demand, and $U_R \leq U_C(x^e, c_1)$, pooling at x^e satisfies the intuitive criterion, including $x^e = 0$.

Proof. Set $u^e = U_C(x^e, c_1)$. Following d , retain types satisfying $u^e \leq \max_{q \in \mathcal{B}(d)} \{qU_C(d, c_1) + (1-q)U_R(c_1)\}$, where $\mathcal{B}$ contains responses optimal under some posterior (Cho and Kreps 1987). Below x^e , neither response improves any type's payoff: use any posterior on survivors and its best response, or the prior if none survive. Above x^e , if concession is possible, monotonicity retains every type; the prior strictly rejects and deters deviations. If concession is impossible, rejection is optimal under every posterior; use a posterior on the surviving types with $U_R = u^e$, or the prior if none survive. On path use the prior and concession. $\square$

C.2. The Closed-Form Region

Assumptions 1–2 imply $\lambda(p - x^*) < z < p - x^*$, $\lambda\bar{c} - a \leq z < \lambda\bar{c}$, and $z \leq \frac{1-p-c_2}{2}$. Thus $L_R = 1$ and $0 < \alpha^* = 1 - \frac{z}{\lambda\bar{c}} \leq \min\{1, \frac{a}{\lambda\bar{c}}\}$: all nonfighters operate after resistance, while some settle after compliance.

C.3. Proof of Proposition 1 and Proposition 2

Proof. Direct substitution in (6) gives the benchmark condition $x(1 - F(p - x)) \leq (p + c_2)[F(p + a) - F(p - x)]$. Under its interiority premises this is $x^2 + (\bar{c} - 2p - c_2)x - (p + c_2)a \leq 0$, with positive root

$$x^\dagger = \frac{-(\bar{c} - 2p - c_2) + \sqrt{(\bar{c} - 2p - c_2)^2 + 4(p + c_2)a}}{2} \quad (7)$$

Binding acceptance gives (2) and $\alpha^\dagger = 1 - F(p - x^\dagger)$.

In either game, Distributional Scope gives strict monotonicity; Assumption 1(i) excludes capped accepted demands. At zero,

$$G(0) = z[L_R - L^C(0)] + (p + c_2 - z)[W_R - W^C(0)]$$

The only action change from resistance to zero compliance is force to settlement, so every posterior accepts zero. If $z < \lambda\bar{c}$, set $t = \max\{p, \frac{z}{\lambda}\}$. The positive-length interval $(t, \min\{\bar{c}, \max\{p + a, \frac{z+a}{\lambda}\}\})$ strictly benefits the target; hence $G(0) > 0$. If $z \geq \lambda\bar{c}$, no type settles after either response, giving $G(0) = 0$ and $x^* = \Omega^* = 0$. This proves Remark 2 without Assumption 2.

In the operation game's closed-form region, substitute $W = \frac{p-x-z}{(1-\lambda)\bar{c}}$, $W_R - W = \frac{x}{(1-\lambda)\bar{c}}$, and $L_R - L = \alpha^*$ into (6). Binding acceptance gives $x^2 + [(1-\lambda)\bar{c} - 2p - c_2 + 2z]x - (1-\lambda)\bar{c}z\alpha^* = 0$, whose positive root is

$$x^* = \frac{-B + \sqrt{B^2 + 4(1-\lambda)\bar{c}z\alpha^*}}{2}, \quad B = (1-\lambda)\bar{c} - 2p - c_2 + 2z \quad (8)$$

Pooling is supported by the preceding belief construction: $U_C = \max\{x, p - c_1, \min\{x + z, 1\} - \lambda c_1\}$ is nondecreasing and dominates $U_R = \max\{-a, p - c_1, z - \lambda c_1\}$; omit operation payoffs in the benchmark. Lemma 5 below proves maximality and outcome uniqueness at the selected demand. $\square$

D. PROOFS OF THE COMPARATIVE STATICS

Write $C = \bar{c}$, $q = 1 - \lambda$, $h = p + c_2$, $X = x^\dagger$, and $Y = x^*$. Derivative subscripts denote partial derivatives.

D.1. Proof of Remark 1

Proof. The binding acceptance condition (6) equates target losses:

$$Y(1 - W(Y)) + hW(Y) + S(Y)(L(Y) - W(Y)) = hW_R + z(L_R - W_R).$$

Therefore

$$\Omega^* = \underbrace{hW_R + z(L_R - W_R)}_{H_R^*} - \underbrace{hW(Y) + S(Y)(L(Y) - W(Y))}_{H_C^*} \quad (9)$$

and subtraction of the benchmark identity gives

$$\Omega^* - \Omega^\dagger = (H_R^* - H_R^\dagger) - (H_C^* - H_C^\dagger). \quad (10)$$

Thus $\Omega^* = 0$ exactly when $H_R^* = H_C^*$. Holding the post-compliance rules fixed in the closed-form region gives $Y_{H_R^*} = \left[1 - W(Y) - \frac{h-z-Y}{qC}\right]^{-1}$. The denominator is positive by decreasing acceptance and below one when $Y < h - z$. $\square$

D.2. Settlement After Compliance

PROPOSITION 7 (Larger Concessions Do Not Increase Settlement While the Operation Is Active): At a common accepted demand x evaluated under the prior, if the operation is active and the operation-settlement cutoff is interior, [Assumption 1\(i\)](#) implies $\kappa_L(x) = \frac{z}{\lambda}$. The cutoff is constant in the demand, and $\Pr^*(\text{settlement} \mid x) = 1 - F\left(\frac{z}{\lambda}\right)$. Settlement becomes less likely as z rises and more likely as λ rises.

Proof. Assumption 1(i) gives $S(x) = z$. Comparing x with $x + z - \lambda c_1$ gives $\kappa_L = \frac{z}{\lambda}$. Lemma 3 then yields $\alpha^* = 1 - F\left(\frac{z}{\lambda}\right)$, $\alpha_z^* = -\frac{f\left(\frac{z}{\lambda}\right)}{\lambda} < 0$, and $\alpha_\lambda^* = f\left(\frac{z}{\lambda}\right) \frac{z}{\lambda^2} > 0$. $\square$

D.3. Gain and Harm

The operational bound limits what force can secure; the stake remains divisible. By contrast, [Spaniel \(2023\)](#) restricts the division of the stake and shows that moderate indivisibility can promote peace by constraining demands or inducing safer offers.

Let the operation deliver fixed gain $g \geq 0$ and impose independently fixed harm $\eta > 0$ throughout the feasible demand set. Maintain interior post-resistance cutoffs and selected demands below the upper bound.

COROLLARY 1 (Gain and Harm): At a common conceded demand x evaluated under the prior, settlement occurs with probability $\alpha(x) = 1 - F\left(\max\left\{p - x, \frac{g}{\lambda}\right\}\right)$. On the active operation branch this is $1 - F\left(\frac{g}{\lambda}\right)$, independent of the target's harm. If $g = 0$ and the operation is active after resistance, the subwar capability raises the accepted demand and coercive ability if and only if $\eta > \lambda(p + c_2)$.

Proof. Settlement beats both alternatives exactly when $c_1 > \max\{p - x, \frac{g}{\lambda}\}$. At $g = 0$ the post-compliance menus coincide. Resistance has operation cutoffs $\frac{p}{q}$ and $\frac{a}{\lambda}$, with activity requiring $qa > \lambda p$. Uniformity gives

$$H_R^* - H_R^\dagger = (qa - \lambda p) \frac{\eta - \lambda h}{\lambda q C}.$$

Since target loss after compliance and $x[1 - F(p - x)]$ strictly increase in x , both comparisons follow. The selected equilibria satisfy the intuitive criterion: $U_C(x, c_1) = \max\{x, p - c_1, x + g - \lambda c_1\}$ increases weakly in x and dominates $U_R(c_1) = \max\{-a, p - c_1, g - \lambda c_1\}$. Apply §C.1. the off-path construction. $\square$

D.4. Proof of Proposition 5

Proof. For any prior, set $Q_0(y) = \int \max\{p - c, y\} dF(c)$ and $Q_1(y) = \int \max\{p - c, y, \min\{y + z, 1\} - \lambda c\} dF(c)$ over $[0, C]$. Then

$$V^* - V^\dagger = [Q_1(Y) - Q_0(Y)] + [Q_0(Y) - Q_0(X)].$$

The first difference is nonnegative. If $Y < X$, every benchmark settler receives strictly less under either benchmark action at Y ; other types weakly lose. Their positive probability makes the second difference negative, without activity or interiority restrictions.

Under uniformity and active interior cutoffs, integrating the optimized payoffs gives

$$V^\dagger = X + \frac{(p - X)^2}{2C}, \quad V^* = Y + \frac{(p - Y - z)^2}{2qC} + \frac{z^2}{2\lambda C}. \quad (11)$$

Specifically, the operation-game integral above Y is $\frac{[k(p - Y - z) - q\frac{k^2}{2} + z\ell - \lambda\frac{\ell^2}{2}]}{C}$, where $k = \frac{p - Y - z}{q}$ and $\ell = \frac{z}{\lambda}$. Subtracting the benchmark value at Y gives $Q_1(Y) - Q_0(Y) = \frac{[\lambda(p - Y) - z]^2}{2\lambda q C}$. Thus payoff falls exactly when the concession loss exceeds this menu gain. $\square$

D.5. Proof of Proposition 6

Proof. Concession loss. Put $b = p - X$. At concession $X - d$, the active war cutoff is $k = \frac{b + d - z}{q}$ and $\ell = \frac{z}{\lambda}$ is constant. Thus $k = b$ at $d_W = z - \lambda b$. After domination, the war cutoff becomes

$b + d$. Hence war falls exactly when $d < d_W$. The cost integral and its derivative on the active branch are

$$K^L(d) = \int_0^k (c + c_2) dF(c) + \lambda \int_k^\ell c dF(c), \quad (\Delta K)' = f(k) \left(k + \frac{c_2}{q} \right) > 0.$$

After domination the derivative is $f(b + d)(b + d + c_2) > 0$. Costs agree at the boundary where the operation interval vanishes. At d_W , identical fighting types incur identical costs, but

$$\Delta K(d_W) = \lambda \int_b^\ell c dF(c) > 0.$$

Continuity and strict monotonicity establish the unique $d_K \in (0, d_W)$ when $\Delta K(0) < 0$. Otherwise $\Delta K(d) > 0$ for $d > 0$. Also $d_W < z$, so $d \geq z$ implies higher cost. These arguments require the stated interior path and positive density, not uniformity.

Operating cost. Under uniformity and active interior cutoffs, put $v = p - x - z$. Integrating gives $K(x, \lambda) = \frac{[v(v+2c_2) + \frac{z^2}{\lambda}]}{2C}$. Consequently $K_x = -\frac{v+c_2}{qC} < 0$, and $K_\lambda < 0$ exactly when

$$\left(\frac{z}{\lambda} \right)^2 > \frac{v(v + 2c_2)}{q^2}. \quad (12)$$

When $c_2 = 0$, activity $\frac{v}{q} < \frac{z}{\lambda}$ ensures this inequality. Differentiating (8) on the strict closed-form branch gives

$$Y_\lambda = \frac{[C(Y - z) + \frac{z^2}{\lambda^2}]}{[2Y + qC - 2p - c_2 + 2z]}. \quad (13)$$

The denominator is positive. The identity $Y > \Omega^* = \frac{Y(h-z)}{qC} + z\alpha^* > z\alpha^*$ implies that the numerator exceeds $q\frac{z^2}{\lambda^2} > 0$. Hence $Y_\lambda > 0$ and $K_\lambda^* = K_\lambda + K_x Y_\lambda < 0$ whenever (12) holds. At $z + a = \lambda C$ these conclusions apply to left derivatives only.

An interval with positive target war costs. Assume $a > 0$, Assumption 1(i), and $C - z > 2p + c_2$. At $\lambda_0 = \frac{z}{C}$, Remark 2 gives $Y = 0$. Assumption 1(ii) and no settlement after resistance hold in a neighborhood. For $p > z$, the positive root immediately above λ_0 has active interior cutoffs. With $u = Y_\lambda$, its right limits are

$$k_0 = \frac{C(p - z)}{C - z} < C = \ell_0, \quad u_0 = \frac{C(C - z)}{C + z - 2p - c_2} > C.$$

The denominator is positive, and $2p + c_2 > 2z$ gives the last inequality. Since

$$2CK_\lambda^* = k^2 - \ell^2 - 2ku + \left(2\frac{c_2}{q}\right)(k - u),$$

its right limit is strictly negative. Continuity gives an interval with $Y_\lambda > 0 > K_\lambda^*$. For $p \leq z$, war is absent nearby and $Y = z(1 - \frac{z}{\lambda C})$, $K^* = \frac{z^2}{2\lambda C}$ give the same signs directly. $\square$

PROPOSITION 8 (The Settlement Cutoff Is the Same Across Equilibria): Consider any equilibrium and any demand conceded with positive probability. If the operation is active after that demand and $S(x) = z$, the coercer settles if and only if $c_1 > \frac{z}{\lambda}$. This cutoff is the same at every such demand and under every off-path belief. If the operation is active and $S(x) = z$ after every conceded demand, so those demands exceed $p - \frac{z}{\lambda}$, then conditional on compliance the settling types are $(\frac{z}{\lambda}, \bar{c}]$ in every equilibrium.

D.6. Proof of Proposition 8

Proof. The coercer's payoffs after compliance, x , $p - c_1$, and $x + z - \lambda c_1$, do not depend on beliefs. Activity means $x > p - \frac{z}{\lambda}$, so settlement occurs exactly above $\frac{z}{\lambda}$. Its probability under posterior μ_x is $1 - \mu_x([0, \frac{z}{\lambda}])$, equaling $1 - F(\frac{z}{\lambda})$ under pooling. A common cutoff does not imply a common posterior probability. If activity fails at another demand, its settlement cutoff is instead $p - x$. $\square$

D.7. Proof of The Resistance-Stage War Comparison

Proof. After resistance, the operation is active exactly when $\frac{p-z}{q} < \frac{z+a}{\lambda}$, equivalently $\lambda(p + a) < z + a$. This is also equivalent to $\frac{p-z}{q} < p + a$, so activity lowers the war cutoff below the benchmark. Assumption 1(ii) gives $\frac{p-z}{q} < C$, making the reduction strict under full support. Domination leaves the benchmark cutoff unchanged. $\square$

D.8. Proof of Proposition 3

Proof. §D.13. Switch accounting gives (19), and (10) proves joint deterioration lowers coercive ability. Substituting (7) and (8) into $\frac{X(C-p+X)}{C}$ and $\frac{Y(qC-p+Y+z)}{qC}$ respectively yields

$$\Omega^\dagger = \frac{h(X+a)}{C}, \quad \Omega^* = \frac{Y(h-z)}{qC} + z\alpha^*.$$

Since $W_R - W(Y) = \frac{Y}{qC}$, this also proves (4). Therefore

$$\Omega^\dagger > \Omega^* \iff \frac{h(X+a)}{C} > \frac{Y(h-z)}{qC} + z\alpha^*. \quad (14)$$

Differentiating the benchmark quadratic gives $X_p = \frac{2X+a}{C-2p-c_2+2X} > 0$ and $\Omega_p^\dagger = \frac{X+a+hX_p}{C} > 0$. Also $(z\alpha^*)_z = 1 - 2\frac{z}{\lambda C}$, so protection peaks at $z = \lambda\frac{C}{2}$. The upper-interval claim is proved below. $\square$

D.9. Proof of Proposition 4

Proof. Set $\beta = C - 2p - c_2 > 0$, $D^\dagger = 2X + \beta$, and $\sigma^j = \Omega_p^j$. Direct differentiation of $X^2 + \beta X = ha$ gives

$$X_p = \frac{2X+a}{D^\dagger}, \quad \sigma^\dagger = \frac{X+a+hX_p}{C}, \quad X_a = \frac{h}{D^\dagger},$$

$$X_{pa} = \frac{\beta^2 + 2\beta h + 2ha}{(D^\dagger)^3} > 0, \quad \sigma_a^\dagger = \frac{1 + X_a + hX_{pa}}{C} > 0.$$

Because $z < \lambda C \leq z + a$, nobody settles after resistance with the operation. Thus Y and σ^* are independent of a . Put $t = p - z$, $s = z(1 - \frac{z}{\lambda C}) > 0$. The four possible regimes give:

- $p < z$: $Y = \Omega^* = s$ and $\sigma^* = 0$.
- $p > z$, no post-compliance war: $Y = \Omega^* = s + \frac{t(t+c_2)}{qC}$ and $\sigma^* = \frac{2t+c_2}{qC} > 0$.
- All three post-compliance actions: $Y^2 + \beta^* Y = qCs$, where $\beta^* = qC - 2p - c_2 + 2z > 0$, $D^* = 2Y + \beta^*$. Then $Y_p = 2\frac{Y}{D^*}$ and $\sigma^* = \frac{Y(qC+c_2+2Y)}{qCD^*} > 0$.
- Dominated operation: $Y^2 + \beta Y = J = Cz + \frac{t(t+c_2)}{q} - ph$. With $D = 2Y + \beta$, domination $p - Y \geq \frac{z}{\lambda}$ implies $J_p = \frac{[\lambda(2p+c_2)-2z]}{q} \geq \frac{\lambda(2Y+c_2)}{q} > 0$. Thus $Y_p = \frac{2Y+J_p}{D}$ and $\sigma^* = \frac{[Y(C+2Y+c_2)+J_p(C-p+2Y)]}{CD} > 0$.

The benchmark algebraic slope extends continuously from zero to infinity as a increases from zero. Its strict increase gives a unique crossing with σ^* . Substituting $a = \frac{X(X+\beta)}{h}$, and putting $k = C - p$, characterizes it by the unique nonnegative root

$$2X_P^3 + 3kX_P^2 + (k^2 + h^2 - 2Ch\sigma^*)X_P - Ch\beta\sigma^* = 0, \quad a_P(p) = X_P \frac{X_P + \beta}{h}. \quad (15)$$

This algebraic extension does not assert equilibria outside $\lambda C - z \leq a < a_{\max} = \frac{p(C-p-c_2)}{h}$. Within that interval, suppression holds exactly when $a > a_P$, with the gap strictly increasing in a . The crossing can lie outside the interval in other action regimes. Adjacent formulas provide one-sided derivatives at boundaries.

Suppression is conditional: $C = 2$, $\lambda = 0.01$, $p = 0.8$, $c_2 = 0.1$, $z = 0.019$, and $a = 0.0011$ satisfy the premises but give $\sigma^\dagger \approx 0.01338 < \sigma^* \approx 0.01862$, with $a_P \approx 0.001547 > a$.

Figure 2. The parameters are $C = 2$, $c_2 = 0.02$, $\lambda = 0.08$, $a = 0.30$. The plotted ranges are $p \in [0.005, 0.90]$ at $z = 0.04$ and $[0.005, 0.74]$ at $z = 0.12$. The premises hold: $p + a < C$, and $X < p$ follows because the convex quadratic $h(p + a) - Cp$ is negative at both endpoints of the larger interval. Throughout, $\sigma^\dagger > \frac{a}{C} = 0.15$. Moreover,

$$(D^\dagger)^2 - (a - \beta)^2 = a[2(C + c_2) - a] > 0, \\ X_{pp} = 2X_p \frac{2 - X_p}{D^\dagger} > 0, \quad \sigma_p^\dagger = \frac{2X_p + hX_{pp}}{C} > 0.$$

Both panels have $s = 0.03$. Post-compliance war begins at $p = z + t_0$, where $t_0 = \frac{[qC - c_2 - \sqrt{(qC - c_2)^2 - 4qCs}]}{2} \approx 0.030852686$. Before that, $\sigma^* \leq \frac{2t_0 + c_2}{qC} < 0.045$. On the three-action branch, $0 < Y_p < 1$, $Y_{pp} = 2Y_p \frac{2 - Y_p}{D^*} > 0$, and $\sigma_p^* = \frac{[2Y_p + (h - z)Y_{pp}]}{qC} > 0$. At $z = 0.04$, this branch ends at

$$p_D = \frac{z}{\lambda} + \frac{E - \sqrt{E^2 - 4qCs}}{2} \approx 0.566203179, \quad E = q\left(C - 2\frac{z}{\lambda}\right) - c_2.$$

Its slope is below the branch formula's value at 0.567, itself below 0.081. At $z = 0.12$, activity persists because $\frac{z}{\lambda} = 1.5 > p$; its endpoint slope is below 0.123.

On panel (a)'s dominated branch, $\beta - J_p = \frac{qC - 2p - c_2 + 2z}{q} > 0$, so $0 < Y_p < 1$. Hence $Y_{pp} = \frac{[2\frac{\lambda}{q} + 2Y_p(2 - Y_p)]}{D} > 0$ and $\sigma_p^* = \frac{2\frac{\lambda}{q} + 2Y_p + hY_{pp}}{C} > 0$. Substitution into the root and slope formulas gives

$$\begin{aligned} p \in [p_D, 0.7] : \quad \sigma^*(p) &\leq \sigma^*(0.7) < 0.172 < 0.394 < \sigma^\dagger(0.566) < \sigma^\dagger(p), \\ p \in [0.7, 0.9] : \quad \sigma^*(p) &\leq \sigma^*(0.9) < 0.508 < 0.546 < \sigma^\dagger(0.7) \leq \sigma^\dagger(p). \end{aligned}$$

Since $Y_p < 1$, cutoffs increase without returning to earlier regimes. These bounds include one-sided derivatives and prove $0 \leq \sigma^* < \sigma^\dagger$ throughout both panels. $\square$

D.10. A Common Interval of Adverse Outcomes

For every fixed p, c_2, z, λ, C admitting the active, interior comparison, all adverse comparisons hold on a nonempty upper interval of admissible audience costs.

Proof. The interval is $a \in [a_{\min}, a_{\max})$, with $a_{\min} = \lambda C - z < a_{\max} = \frac{p(C - p - c_2)}{h} \leq C - p$. Here Y is fixed while X increases toward p . Benchmark punishment increases toward p , but $H_R^* = z + (p - z)\frac{h - z}{qC} < p$, since $qC > 2p + c_2$. Benchmark compliance harm falls toward zero while $H_C^* > 0$. Thus joint deterioration holds on an upper interval. Also $Y < p - z$ implies $X - Y > z$ near the endpoint. Proposition 9 then gives every-type weak payoff loss, strict benchmark-settler loss, and strictly greater war and cost.

To prove $a_P < a_{\max}$, normalize p, c_2, z, a, X, Y by C and relabel c_2 as c . This preserves power slopes. Put $h = p + c$, $t = p - z$, $B = \frac{c}{h}$. Feasibility gives $0 < Y < t < p$, $q > 2p + c$, $q > t + B$, and $h < 1$. Taking the power derivative before the limit $a \rightarrow a_{\max}$ gives

$$L = \lim_{a \rightarrow a_{\max}} \sigma^\dagger = \frac{p}{h} + \frac{p(1 + h)}{1 - c}, \quad \sigma^* = \frac{Y(q + c + 2Y)}{q(q - 2t - c + 2Y)}.$$

Set $D = q - 2t - c > 0$. The numerator determining σ_Y^* is $4Y^2 + 4DY + (q + c)D > 0$. Therefore

$$\sigma^* < J(t, q), \quad J(t, q) = \frac{t(q + c + 2t)}{q(q - c)}.$$

The numerator determining J_q is $-t[q^2 + (2q - c)(c + 2t)] < 0$, so

$$\sigma^* < J(t, \max\{2p + c, t + B\}) \leq J(p, \max\{2p + c, p + B\}) =: R.$$

The last inequality follows from increasing J in t at fixed second argument. Where that argument is $t + B$, the derivative numerator is $(5B - 4c)t^2 + 6B(B - c)t + (B + c)B(B - c) \geq 0$ because $B \geq c$. For $c = 0$ only the constant branch applies.

If $h^2 \geq c$, $R = 1$ and $L - 1 = (1 + p)\frac{h^2 - c}{h(1 - c)} \geq 0$. If $h^2 < c$,

$$L - R = \frac{p(1 - h)(c - h^2)[cp(3 - c) + c + p^2(p + 2)]}{(1 - c)h(c(1 - c) + p^2)(c(p + 1) + p^2)} > 0.$$

Thus $L > \sigma^*$. Continuity and strict increase of $\sigma^\dagger$ imply $a_P < a_{\max}$. Intersecting the upper intervals proves the claim. Strict comparisons persist under nearby primitives within the same strict regime. $\square$

The result does not bound the fraction of admissible audience costs in this interval. For $C = 2$, $c_2 = 0.02$, $\lambda = 0.08$, $p = 0.30$, and $z = 0.04$, all adverse comparisons hold for $a > 0.280$ within $[0.12, 1.575)$. Exact thresholds vary with the military fundamentals.

D.11. Payoff Loss and War

Let $b = p - X$ and $d = X - Y$.

PROPOSITION 9 (Coercer Payoff Loss and Increased War Share a Threshold): Suppose $S(x^*) = z$ and $b = p - x^\dagger \in (0, \bar{c})$. Every coercer type is weakly worse off with the operation if and only if (16) holds. The same condition holds if and only if war after compliance is weakly more likely. War is strictly more likely exactly when the inequality is strict. If the condition holds and $d > 0$, every benchmark settler strictly loses, expected joint conflict cost is strictly higher, and coercive ability is strictly lower. If $d = 0$ and the condition holds, payoffs, war, conflict cost, and coercive ability are unchanged.

$$\underbrace{d}_{\text{concession lost}} \geq \underbrace{\max\{0, z - \lambda b\}}_{\text{net gain at the war margin}}. \quad (16)$$

Proof. This comparison allows any atomless full-support prior on $[0, C]$. Payoffs are $U^\dagger(c) = \max\{p - c, X\}$ and $U^*(c) = \max\{p - c, Y, Y + z - \lambda c\}$. Benchmark settlers require $d \geq 0$ for every-type weak loss. Given this, only the operation can improve payoffs. Its advantage over

$U^\dagger$ increases with slope q below b and decreases with slope $-\lambda$ above it. Its maximum is $z - \lambda b - d$, proving the payoff equivalence.

The actual operation-game war cutoff is $k = \min\left\{b + d, \frac{b+d-z}{q}\right\}$, including domination and, through the clamped CDF, exterior cutoffs. Because b is interior and F strictly increases there, $F(k) \geq F(b)$ exactly when both $d \geq 0$ and $d \geq z - \lambda b$. Strict war requires both strict inequalities, equivalently strict (16).

With $d > 0$, every $c > b$ obtains less than X from war and settlement, and at most $X - \lambda(c - b) < X$ from the operation. Types below b fight in both games. If $k > b$, additional war adds positive costs. If $k = b$, $d = z - \lambda b > 0$ and types in $(b, \min\{\frac{z}{\lambda}, C\})$ instead operate, also adding positive costs. Finally $\Omega^* = Y[1 - F(k)] \leq Y[1 - F(b)] < X[1 - F(b)] = \Omega^\dagger$. With $d = 0$, the operation never improves the benchmark menu; the maintained tie rule preserves all outcomes.

In the closed-form region, $d > z$ is a stronger sufficient test, equivalent to

$$ha + z(2p + c_2 - qC - z - qC\alpha^*) > \lambda CX. \quad (17)$$

Indeed, $\varphi(t) = t^2 + (qC - 2p - c_2 + 2z)t - qCz\alpha^*$ has positive root Y and $\varphi(-z) = z(2p + c_2 - qC - z - qC\alpha^*) < 0$. As $X - z \geq -z$, $X - Y > z$ exactly when $\varphi(X - z) > 0$. Substituting the benchmark quadratic yields the displayed inequality, without evaluating an off-branch probability. $\square$

D.12. An Unused Operation

Under Assumption 1, suppose a positive measure operates after resistance but the operation is dominated after the selected concession. Then threats and assurances weaken, concessions and coercive ability fall, every coercer type weakly loses, and war and cost rise, without additional interiority restrictions.

Proof. Domination gives $z \leq \lambda(p - Y) < \lambda C$. Remark 2 implies $Y > 0$, so $\frac{z}{\lambda} < p \leq h$. Let $m_W > 0$ and $m_S \geq 0$ be the masses switching from war and settlement after resistance. If $p + a < C$, uniformity gives $m_W = \frac{z+qa-\lambda p}{qC}$ and $m_S \leq q\frac{m_W}{\lambda}$. Hence

$$H_R^* - H_R^\dagger = zm_S - (h - z)m_W \leq \left(\frac{z}{\lambda} - h\right)m_W < 0.$$

If $p + a \geq C$, $m_S = 0$ gives the same strict sign. At Y both games' optimized compliance menus equal $\max\{Y, p - c_1\}$. The benchmark acceptance margin is therefore $H_R^\dagger - H_R^* > 0$, implying $X > Y$. Every type weakly loses and benchmark settlers strictly lose with the operation. Since $0 < p - Y < C$, full support gives $F(p - Y) > F(p - X)$ even if $p - X \leq 0$. Additional war strictly raises post-compliance harm and total cost. The strictly increasing function $x[1 - F(p - x)]$ gives lower coercive ability. No positive-probability type operates after compliance. $\square$

D.13. Interior Switch Accounting

At fixed history (y, r) with $S(y) = z$ and interior cutoffs, put $M(y) = z + qa\mathbb{1}[r = 1] - \lambda(p - y)$ and $M_R = M(0)$. Activity means $M(y) > 0$. Uniformity gives switching masses $m_W = \frac{M(y)}{qC}$ and $m_S = \frac{M(y)}{\lambda C}$ from the cutoff gaps. Excluding the retained transfer, harm changes by

$$zm_S - (h - z)m_W = M(y)\frac{z - \lambda h}{\lambda q C}. \quad (18)$$

Total target loss also changes by ym_W . At resistance $y = 0$. Since $M_R - M(x) = qa - \lambda x$, the difference between resistance and compliance harm changes has sign $(z - \lambda h)(qa - \lambda x)$.

In the closed-form region, settlement-to-operation switching after resistance is truncated at C . There $m_S = \frac{C - p - a}{C}$. After compliance, adjusting from X to Y adds $\frac{h(X - Y)}{C}$ to harm. Thus joint deterioration holds exactly when

$$q(C - p - a)z < M_R(h - z), \quad M(Y)\frac{z - \lambda h}{\lambda q} + h(X - Y) > 0. \quad (19)$$

E. PURE TARGET RESPONSES AND D1

Write $\ell = \frac{z}{\lambda}$, $s_R = \max\{p + a, \frac{z + a}{\lambda}\}$, and $\Phi(d) = \int_0^{\bar{c}} [u_2^C(d, c_1) - u_{2,R}(c_1)] dc_1$, using optimal continuations. Let $x^e = \max\{d \in [0, 1] : \Phi(d) \geq 0\}$ and

$$\underline{x} = \begin{cases} 0 & \text{benchmark: } \bar{c} > p + a \\ p + c_2 & \text{benchmark: } \bar{c} \leq p + a \\ 0 & \text{operation: } \bar{c} > s_R \text{ or } z > \lambda p \\ \min\{z, p + c_2 - z\} & \text{operation: } \bar{c} \leq s_R \text{ and } z \leq \lambda p \end{cases}$$

LEMMA 5 (Equilibrium Outcomes Under Pure Target Responses): Under Assumption 1 and the maintained tie rules, the largest accepted demands in pure-target-response perfect Bayesian equilibria are exactly $[\underline{x}, x^e]$, also under the intuitive criterion. Each admits a pool reproducing every type's military action and both payoffs. Only neutral types with $U_C(x, c_1) = U_R(c_1)$ can use other signals. Randomizing coercer demands adds no terminal outcomes.

Proof. Partition types by resistance action and preferred nonwar compliance action. For pairs (W, O), (O, O), (W, S), (O, S), (S, S), respectively, the target's nonwar surplus is $J - d$, with $J = p + c_2 - z, 0, p + c_2, z, 0$; compliance war gives zero. Keep nonempty regions, including singleton types. Nonwar compliance is optimal by $d = J$ in each region: switching thresholds against war are at most $p - z$ or p ; following nonwar resistance it is already optimal at zero. Therefore all posteriors accept $d \leq \min J$, and every larger demand admits strict rejection. Assumption 1(i) strictly excludes capped demands.

Resistance settlement exists exactly when $\bar{c} > s_R$; operation after both responses to zero exists exactly when $z > \lambda p$. Otherwise $\ell \leq p \leq b_R = \frac{p-z}{1-\lambda} < \bar{c}$, giving the displayed operation floor. Benchmark resistance settlement exists exactly when $\bar{c} > p + a$, giving its floor. The strict monotonicity proved above and rejection at the demand ceiling imply that the prior accepts exactly $[0, x^e]$, with $\Phi(x^e) = 0$.

Zero must be accepted. An accepted set without a maximum would leave a positive mass of high-cost nonfighters with an unattained payoff supremum, also unattainable by mixing. At its maximum x , every type obtains $U_C(x, c_1)$ by monotonicity and Lemma 2. If $x > 0$, strict-gain types uniquely choose x ; others always fight and contribute zero acceptance surplus. Removing or reweighting them cannot change its sign, so Bayes' rule requires $\Phi(x) \geq 0$. At zero this already holds. Hence $\underline{x} \leq x \leq x^e$.

Conversely, pool at any such x , accept $d \leq x$, and reject $d > x$. Below x , equilibrium-dominance survivors are neutral fighters, if any; their posterior supports concession at indifference. Otherwise use the prior. Above x , if concession is possible under some posterior, monotonicity retains every type, permitting the strictly rejecting posterior constructed above. If concession is impossible, use neutral survivors or, if none, any posterior. Rejection never improves coercer payoffs. This proves intuitive-criterion sufficiency, including zero and capped deviations.

Only neutral types can change signals or mix; their actions and both payoffs remain unchanged. At $x > 0$ they always fight; at zero they can also operate after both responses. Bayes-consistent reweighting therefore preserves unconditional outcomes. Used rejected signals require strict target rejection. If neutral types exist, the criterion forces acceptance below x ; otherwise only the universal floor binds. $\square$

E.1. Zero Demands and Boundary Cases

Zero is attainable exactly when the displayed floor is zero; the strict endpoints follow from continuation tie-breaking. Benchmark $\bar{c} \leq p + a$ instead yields the singleton $\{p + c_2\}$, since every type imposes target loss $p + c_2$ after either response at that demand. In the operation game $x^e = 0$ exactly when $\ell \geq \bar{c}$, as proved above; then everyone can instead demand one and be rejected.

The positive-floor operation interval collapses exactly when $z = \lambda p$ and $p + c_2 = 2z$. If $z < \lambda p$, types in (p, b_R) switch from resistance war to compliance settlement, contributing strictly positive surplus at the floor. If $z = \lambda p$, then $b_R = \ell = p$; $p + c_2 < 2z$ gives positive operation-to-settlement surplus above p , and $p + c_2 > 2z$ positive war-to-operation surplus below p . At equality every contribution is zero.

E.2. Minimum War and Conflict Cost

Using clipping $[v]_0^{\bar{c}} = \min\{\bar{c}, \max\{0, v\}\}$, set $\kappa(x) = [\min\{p - x, \frac{p-x-z}{1-\lambda}\}]_0^{\bar{c}}$ and $\ell_x = [\max\{p - x, \ell\}]_0^{\bar{c}}$. In the benchmark both equal $[p - x]_0^{\bar{c}}$.

COROLLARY 2 (The Largest Demand Minimizes Conflict): Under Lemma 5, x^e maximizes every coercer type's payoff and unconditional retained concession and minimizes unconditional war and joint conflict cost. Minima need not be unique.

Proof. Integrating gives $W(x) = \frac{\kappa(x)}{\bar{c}}$ and $K(x) = \frac{[(1-\lambda)\kappa(x)^2 + \lambda\ell_x^2 + 2c_2\kappa(x)]}{2\bar{c}}$. Both cutoffs decrease weakly; thus W, K decrease and $x(1 - W)$ and type payoffs increase weakly. Actual resistance never ends in settlement because $a > 0$ makes available compliance preferable, so joint surplus is $1 - K$. Lemma 5 extends these quantities to every pure-response equilibrium. $\square$

Selected increases raise the minimum attainable war or cost, without ranking every cross-game equilibrium pair. Conditional outcomes can differ when neutral types use rejected signals; they are undefined if compliance never occurs.

E.3. D1 and the Selected Pooling Equilibria

Let $\mathcal{D}_c$ and $\mathcal{D}_c^0$ contain concession probabilities making a deviation strictly profitable or payoff-equivalent. D1 eliminates c if $\mathcal{D}_c \cup \mathcal{D}_c^0$ is a proper subset of another type's strict-gain set.

PROPOSITION 10 (Positive Concessions Fail D1 Under Pure Target Responses): Under Assumption 1 and concession at indifference, every pure-target-response perfect Bayesian equilibrium with largest accepted demand $\bar{x} > 0$ gives type payoff $U_C(\bar{x}, c_1)$. If its compliance war cutoff is interior, it fails D1. For $c_2 > 0$, failure is independent of the tie rule.

E.4. Proof of Proposition 10

Proof. The maximum/payoff argument in Lemma 5 uses only optimality, so applies without its refinement. Let $k = \kappa(\bar{x}) > 0$ be untruncated. The war cutoff is continuous, strictly decreasing, and below the resistance cutoff. Choose $d > \bar{x}$ with $0 < k' = \kappa(d) < k$, $d < p$ in the benchmark or $d + z < p$ with the operation. Such a deviation exists on both active and dominated branches.

Write $E = U_C(\bar{x}, c) - U_R(c)$ and $\Delta = U_C(d, c) - U_R(c)$. Profit requires $q\Delta > E$. For $c \leq k'$, $(\mathcal{D}, \mathcal{D}^0) = (\emptyset, [0, 1])$; for $k' < c \leq k$, it is $((0, 1], \{0\})$. Above k , the weak-gain set $[\frac{E}{\Delta}, 1]$, with

$\frac{E}{\Delta} > 0$, is a proper subset of the switchers' strict-gain set. D1 therefore retains exactly $[0, k]$. These types always fight after resistance; after compliance they either still fight or impose loss at most $d + z < p \leq p + c_2$ (benchmark: $d < p$). Every admissible belief consequently induces concession, and switchers profit. An on-path d would be rejected, but only neutral types could use it, also restricting Bayes' support to $[0, k]$ and contradicting rejection.

For $c_2 > 0$, instead choose $d \in (p - z, p + c_2 - z)$, or $(p, p + c_2)$ in the benchmark. All surviving types then leave war and strictly benefit the target, removing reliance on tie-breaking. When $c_2 = 0$, off-path rejection at indifference could instead be supported by always-fighting types.

Positivity is necessary: take $(p, c_2, a, \lambda, z, \bar{c}) = (.55, .02, .30, .08, .12, 2)$. Pooling at zero gives target surplus .03 and an interior war cutoff. After every positive demand, put belief on $c = 1$: this type operates after either response to zero, earns .04, and has strict-gain set $(0, 1]$, so D1 cannot eliminate it. Its compliance loss is $d + .12$ up to .88, one up to .92, and d thereafter, always exceeding resistance loss .12. Strict rejection deters all deviations and also satisfies the intuitive criterion. With $z = .04$ instead, every posterior accepts $0 < d < .04$: fighters give zero surplus and settlers avoid resistance loss .57 or .04, precluding a zero equilibrium.

Interiority is necessary too: benchmark $p + a \geq \bar{c}$ permits a D1 equilibrium accepting through $p + c_2$. Above that demand, a D1-admissible belief on $c = 0$ strictly favors resistance. $\square$

F. EXTENSIONS AND SCOPE

F.1. Escalation Risk

An operation escalates to war with probability $\xi \in [0, 1 - \lambda)$, incurring λc_1 before any additional war cost. After transfer y , its payoffs are

$$\begin{aligned} U_1^L(y, c_1; \xi) &= (1 - \xi)(y + S(y)) + \xi p - (\lambda + \xi)c_1, \\ U_2^L(y; \xi) &= (1 - \xi)(1 - y - S(y)) + \xi(1 - p - c_2). \end{aligned} \tag{20}$$

Comparison with war and settlement gives the active-operation cutoffs

$$\begin{aligned}\kappa_W(y; \xi) &= (1 - \xi) \frac{p - y - S(y)}{1 - \lambda - \xi}, \\ \kappa_L(y, r; \xi) &= \frac{(1 - \xi)S(y) + \xi(p - y) + a\mathbb{1}[r = 1]}{\lambda + \xi}.\end{aligned}\tag{21}$$

For uncapped concessions, $\partial \frac{\kappa_L}{\partial} y = -\frac{\xi}{\lambda + \xi}$: positive escalation risk makes larger concessions restrain operation use.

Let W_ξ, L_ξ denote probabilities of direct war and either military action. They equal $F(\kappa_W), F(\kappa_L)$ on active branches and both equal $F(p - y + a\mathbb{1}[r = 1])$ when the operation is dominated. Subscript R evaluates $(y, r) = (0, 1)$; superscript C evaluates $(x, 0)$. Total war probabilities and target losses are

$$\begin{aligned}R_{\xi, R} &= W_{\xi, R} + \xi(L_{\xi, R} - W_{\xi, R}), \\ R_\xi^C(x) &= W_\xi^C(x) + \xi(L_\xi^C(x) - W_\xi^C(x)).\end{aligned}\tag{22}$$

$$\begin{aligned}D_{\xi, R} &= (p + c_2)R_{\xi, R} + (1 - \xi)z(L_{\xi, R} - W_{\xi, R}), \\ D_\xi^C(x) &= x[1 - R_\xi^C(x)] + (p + c_2)R_\xi^C(x) + (1 - \xi)S(x)[L_\xi^C(x) - W_\xi^C(x)].\end{aligned}\tag{23}$$

Thus

$$x_\xi^* = \max\{x \in [0, 1] : D_\xi^C(x) \leq D_{\xi, R}\}, \quad \Omega_\xi^* = x_\xi^*[1 - R_\xi^C(x_\xi^*)].\tag{24}$$

For active operations after both responses, interior war cutoffs, an interior compliance operation–settlement cutoff, $L_{\xi, R} = 1$, and $S(x) = z$, define

$$\begin{aligned}A_\xi &= 1 - \lambda - \xi, \quad d_\xi = \lambda + \xi, \quad t_\xi = 1 - \xi, \\ w_\xi &= \frac{t_\xi(p - z)}{A_\xi \bar{c}}, \quad \ell_\xi = \frac{t_\xi z + \xi p}{d_\xi \bar{c}}, \\ r_\xi &= \xi \ell_\xi + t_\xi w_\xi, \beta_\xi = \frac{\xi^2}{d_\xi \bar{c}} + \frac{t_\xi^2}{A_\xi \bar{c}}, \\ \eta_\xi &= \frac{t_\xi}{A_\xi \bar{c}} - \frac{\xi}{d_\xi \bar{c}}.\end{aligned}\tag{25}$$

Then $W_\xi^C = w_\xi - t_\xi \frac{x}{A_\xi \bar{c}}$, $L_\xi^C = \ell_\xi - \xi \frac{x}{d_\xi \bar{c}}$, and $R_\xi^C = r_\xi - \beta_\xi x$. Substitution gives

$$\begin{aligned}
H_\xi &= \xi(p + c_2) + t_\xi(p + c_2)w_\xi + t_\xi z(1 - w_\xi), \\
C_\xi &= H_\xi - [(p + c_2)r_\xi + t_\xi z(\ell_\xi - w_\xi)],
\end{aligned} \tag{26}$$

$$\begin{aligned}
B_\xi &= 1 - r_\xi - (p + c_2)\beta_\xi + t_\xi z\eta_\xi. \\
\beta_\xi x^2 + B_\xi x - C_\xi &= 0.
\end{aligned} \tag{27}$$

For $C_\xi > 0$ and decreasing acceptance margin, the positive root is

$$x_\xi^* = \frac{-B_\xi + \sqrt{B_\xi^2 + 4\beta_\xi C_\xi}}{2\beta_\xi}. \tag{28}$$

At $\xi = 0$, this recovers (8).

PROPOSITION 11 (Coercive Losses and Increased Conflict Survive Small Escalation Risks): Suppose the operation-game closed-form and active, interior conditions hold strictly at $\xi = 0$. For some $\bar{\xi} > 0$, the selected pooling equilibrium exists continuously on $[0, \bar{\xi})$. Strict $\Omega^* < \Omega^\dagger$ and $K^* > K^\dagger$ comparisons persist, as does $\frac{\partial K^*}{\partial \lambda} < 0$. If $x^\dagger > x^* + z$, every type remains weakly worse off, a positive measure strictly worse off, and both direct and total war remain strictly more likely than in the benchmark.

Proof. Write $G_\xi = D_{\xi,R} - D_\xi^C$. Assumption 1(ii) gives $G_{0'}(x^*) < 0$. Strict regime conditions make G_ξ smooth near $(x^*, 0)$, so the implicit function theorem supplies a locally unique, smooth root. Proposition 2 gives $G_0 < 0$ on $[x^* + \delta, 1]$ for every sufficiently small $\delta > 0$. Continuity of optimal-action probabilities gives uniform convergence there. Thus higher demands remain rejected, and the root remains globally largest. Each action's payoff weakly increases with compliance and the demand, so [SI §C.1. the pooling construction](#) applies.

The cutoffs and Ω_ξ^* are continuous. Under $x^\dagger > x^* + z$, continuity preserves $x^\dagger > x_\xi^* + z$. The operation payoff is $(1 - \xi)(x_\xi^* + z) + \xi(p - c_1) - \lambda c_1 \leq U^\dagger(c_1)$. War is unchanged and settlement lower, establishing every-type weak loss and strict loss for common settlers. Proposition 9's strict direct-war comparison persists by continuity; total war is at least direct war.

Expected cost is

$$K_\xi^* = \int_0^{\kappa_W} (c_1 + c_2) dF(c_1) + \int_{\kappa_W}^{\kappa_L} [\lambda c_1 + \xi(c_1 + c_2)] dF(c_1),$$

with cutoffs evaluated at x_ξ^* . Smoothness of the root, cutoffs, and integrands makes both K_ξ^* and $\frac{\partial K_\xi^*}{\partial \lambda}$ continuous. Their strict comparisons persist. Intersecting these neighborhoods proves the claim. $\square$ The result is local and uses strict regime conditions. Equality in the baseline payoff threshold alone does not establish escalation robustness.

F.2. Power Shift and Partial Recontestation

Maintain interior benchmark war cutoffs, active operations with interior compliance cutoffs, $L_R = 1$, and $S(x) = z$ at selected demands in the shifted game.

REMARK 4 (Power Shift): Let $\gamma \in [0, 1 - p]$. War after compliance pays $p + \gamma x - c_1$ and $1 - p - c_2 - \gamma x$. Reopening share ρ of the transfer gives $\gamma = (1 - \rho)(1 - p)$. Resistance is unchanged. Compliance cutoffs are $\kappa^\dagger(x) = p - (1 - \gamma)x$, $\kappa_W(x) = \frac{p - (1 - \gamma)x - z}{1 - \lambda}$, and $\kappa_L = \frac{z}{\lambda}$. Thus $\alpha^* = 1 - F(\frac{z}{\lambda})$ remains constant. Effective payment $P = x[1 - (1 - \gamma)W]$ equals $H_R^* - H_C^*$, with

$$P^\dagger = \frac{(p + c_2)((1 - \gamma)x^\dagger + a)}{\bar{c}}, \quad P^* = \frac{(1 - \gamma)x^*(p + c_2 - z)}{(1 - \lambda)\bar{c}} + z\alpha^*.$$

Demands are the positive roots below. Coercive ability remains $\Omega = x(1 - W)$ in each game; $\gamma = 0$ recovers the baseline.

$$\begin{aligned} (1 - \gamma)^2 x^2 + (\bar{c} - (1 - \gamma)(2p + c_2))x - (p + c_2)a &= 0, \\ (1 - \gamma)^2 x^2 + ((1 - \lambda)\bar{c} - (1 - \gamma)(2p + c_2) + 2(1 - \gamma)z)x - (1 - \lambda)\bar{c}z\alpha^* &= 0. \end{aligned} \tag{29}$$

Payoff intersections give the cutoffs. Target compliance loss is $P + (p + c_2)W + z(L - W)$. Equating this to resistance loss, with $W_R - W = (1 - \gamma)\frac{x}{(1 - \lambda)\bar{c}}$, gives the operation quadratic; omitting the operation gives the benchmark quadratic. Assumption 1(ii) remains sufficient for monotonicity. Each action's payoff weakly increases with compliance and demand, permitting [SI §C.1. the pooling construction](#). For $\gamma > 0$, all types strictly lose from smaller demands, which

therefore impose no intuitive-criterion belief restriction. Outside the stated regime, use the globally optimal menu.

REMARK 5 (Payoffs and War under a Power Shift): On this active, interior branch, every-type weak loss holds exactly when $x^\dagger \geq x^*$ and (30) holds. War is strictly more likely if and only if (31) holds. For $\gamma > 0$ and $x^\dagger > x^*$, the war condition is stronger, common fighters lose $\gamma(x^\dagger - x^*)$, and strict (30) implies almost-everywhere strict payoff loss. At $\gamma = 0$, the conditions recover Proposition 9's weak payoff and strict-war comparisons, and common fighters are indifferent. The stronger sufficient condition $(1 - \gamma)(x^\dagger - x^*) > z$ ensures both comparisons.

$$x^\dagger + \lambda \kappa^\dagger(x^\dagger) \geq x^* + z. \quad (30)$$

$$(1 - \gamma)(x^\dagger - x^*) + \lambda \kappa^\dagger(x^\dagger) > z. \quad (31)$$

Common settlers require $x^\dagger \geq x^*$. Given this, benchmark war and settlement dominate their operation-game counterparts. The benchmark payoff advantage over the operation decreases until $\kappa^\dagger(x^\dagger)$ and increases afterward. Its minimum is $x^\dagger + \lambda \kappa^\dagger(x^\dagger) - x^* - z$, proving (30). Comparing war cutoffs gives (31).

F.3. The Operation's Audience Cost

REMARK 6 (An Audience Cost on the Operation): Let the operation incur $a_L \in [0, a]$ after either response. Resistance cutoffs become $\kappa_{L,R} = \frac{z+a-a_L}{\lambda}$ and $\kappa_{W,R} = \frac{p-z+a_L}{1-\lambda}$. Positive operation use requires $\max\{0, \kappa_{W,R}\} < \min\{\bar{c}, \kappa_{L,R}\}$. For interior benchmark resistance cutoff, this and strict war reduction are equivalent to $z + a - a_L > \lambda(p + a)$.

After uncapped compliance, $\kappa_L = \frac{z-a_L}{\lambda}$ and $\kappa_W = \frac{p-x-z+a_L}{1-\lambda}$. Under Assumption 1(ii), activity requires exactly $a_L < z$ and $a_L < z - \lambda(p - x)$. Then $\alpha = 1 - F(\kappa_L)$; otherwise $\alpha = 1 - F(p - x)$.

Maintain full reach, interior war and compliance operation-settlement cutoffs, active operations, and $L_R = 1$, requiring $a_L \leq z + a - \lambda\bar{c}$. The selected demand solves (32). Both x^* and Ω^* increase in a_L when $z < p + c_2$, with one-sided comparisons at boundaries.

$$x^2 + x((1 - \lambda)\bar{c} - 2p - c_2 + 2z - a_L) - (1 - \lambda)\bar{c}z\alpha^* = 0. \quad (32)$$

Payoff intersections establish the cutoff and activity conditions. Since $W_R - W = \frac{x}{(1-\lambda)\bar{c}}$, target indifference gives (32) and $\Omega^* = x^* \frac{p+c_2-z}{(1-\lambda)\bar{c}} + z\alpha^*$. Differentiation yields

$$\frac{\partial x^*}{\partial a_L} = \frac{x^* + (1 - \lambda)\frac{z}{\lambda}}{2x^* + (1 - \lambda)\bar{c} - 2p - c_2 + 2z - a_L} > 0.$$

Compliance-loss derivatives on active, dominated, and no-war branches are respectively $1 - \frac{2p+c_2-2x-2z+a_L}{(1-\lambda)\bar{c}}$, $1 - \frac{2p+c_2-2x}{\bar{c}}$, and 1, all positive by Assumption 1(ii) and $a_L < z$. Loss is continuous across branches. Capped demands are rejected: $1 - a_L > p$ excludes war, and compliance costs at least $x > 1 - z \geq \max\{p + c_2, z\}$, bounding resistance loss. The root is therefore globally largest. Action payoffs weakly increase with compliance and demand, so SI §C.1. the pooling construction applies. These derivatives establish the stated branch comparison, without a global ordering after operation dominance.