

Brook no compromise: How to negotiate a united front

Elaine Yao (Princeton)

April 1, 2025

- ▶ How do groups overcome their internal differences to get things done?
- ▶ Political parties, international organizations, and rebel groups all want to form a united front – but need to decide **on whose terms**
- ▶ How do groups resolve **the tension between cooperation and conflict**?
- ▶ In successful united fronts, who gets their way? Why do groups sometimes fail to unite?

- ▶ Can moderates and the Freedom Caucus within House GOP unite to repeal the ACA?
- ▶ Moderates preempt FC by scheduling a vote on their preferred bill
- ▶ FC withholds support; “we’re going to be living with Obamacare for the foreseeable future”

How a secret Freedom Caucus pact brought down Obamacare repeal



Freedom Caucus vice chairman Jim Jordan hatched the pact to bind the Freedom Caucus together in negotiations and ensure the White House or House leaders could not peel them off one by one. | AP Photo

- ▶ Miscoordination happens even when coordination seems like it would be in the interest of both sides:
 - ▶ Failed Islamist–Communist alliances in anticolonial independence movements
 - ▶ Conflicts between “moral hazard hawks” and “liquidity doves” over bailouts
- ▶ ... but is not inevitable: House GOP does sometimes pass legislation with the Freedom Caucus’ support ([Green 2019](#)), bailout deals do pass!
- ▶ Each situation has unique features which generate variation and discretion over who acts first, how, and the consequences
- ▶ **What kinds of information, uncertainty, and action matter? Can we use these to generate systematic explanations about coordination outcomes?**

Argument

- ▶ The process of coordination involves gauging your counterpart's **willingness to compromise** and **opportunities to play hardball**
 - ▶ Willingness to compromise: Would my opponent support my cause over the status quo, if they had no other choice?
 - ▶ Hardball: An action which forces my opponent to choose between my cause and the status quo
- ▶ The **threat of hardball** leads groups to **learn about each other's willingness to compromise** over time. This determines **the timing of when** groups use hardball tactics and in turn, whether a united front is formed

Preview of model

- ▶ Dynamic two-player coordination game with uncertainty over willingness to compromise and opportunities for “hardball”
- ▶ In the unique PBE, players who are willing to compromise delay in order to learn about their opponents
- ▶ Factors which inhibit learning extend delay – Potentially bad for the individual player; but mean that avoidable miscoordination is less likely

Related literature: Bargaining

- ▶ Legislative bargaining (Baron and Ferejohn 1989; Baron 1991; Romer and Rosenthal 1978; Krehbiel 1996, 1998; Banks and Duggan 2000, 2006)
 - **political frictions** in the bargaining protocol generated by “strong” institutional arrangements
- ▶ Reputational bargaining (Abreu and Gul 2000; Fudenberg and Tirole 1986; Milgrom and Roberts, 1982; Reich 2024)
 - **reputational frictions** generate by incentives to hold out for high demands (pose as a “behavioral type”)

Relationship to literature

- ▶ I capture **reputational frictions** through uncertainty over willingness to compromise, and **political frictions** through arrival of hardball opportunities
 - ▶ Learning about reputation imperfect bc of political frictions
 - ▶ “Hardball actions” flip screening intuition of reputational bargaining, and generate costs to delay even with no discounting
 - ▶ Political frictions discipline multiplicity in coordination games
- ▶ Dynamic trade-off between preemption and caution parallels **preemption games with private info** [Hopenhayn & Squintani \(2011\)](#), [Weeds \(2002\)](#), [Fudenberg & Tirole \(1985\)](#), [Bobtcheff, Levy and Mariotti \(2022\)](#), [Shahanaghi \(2024\)](#)
 - ▶ Behavior (e.g. filing a patent) is publicly observable, but payoff-relevant characteristics (e.g. quality of innovation) are privately known

Model setup

Policies, players, and types

- ▶ Three **policies**: SQ , A , and B
- ▶ Two **players**, a and b (“groups”)
 - ▶ a likes A best, b likes B best
- ▶ A “**hard**” **type** prefers SQ to the other group’s policy; a “**soft**” **type** prefers the other group’s policy over SQ
 - ▶ If player a is a **hard type**, $A \succ_a SQ \succ_a B$
 - ▶ If player a is a **soft type**, $A \succ_a B \succ_a SQ$
- ▶ Each player i holds a prior p_i that its opponent is a hard type

Actions & Timing

Actions:

- ▶ Hardball: Take-it-or-leave-it offer between one policy and status quo
 - ▶ If the opponent accepts, the policy is implemented
 - ▶ If the opponent rejects, status quo stays in place

Timing:

- ▶ Time is continuous and infinite
- ▶ Groups can only play hardball when they (privately) receive an **opportunity**
- ▶ *Interpretation:* Opportunities capture frictions like procedural constraints, group discipline, resolve, or resource constraints
- ▶ Opportunities arrive $\sim \text{Poisson}(\mu)$ where $\mu_{hard} \geq \mu_{soft}$

Timing

1. **Before the game:** SQ is in place.
 2. **Once game starts:** Players receive Poisson opportunities to act
 - 2a. If they pass on an opportunity, it is unobserved, game continues
 - 2b. If they use the opportunity to commit to a policy, game “ends”
 - ▶ If opponent **rejects**, get SQ forever
 - ▶ If opponent **accepts**, new policy is instated permanently
- ▶ Solution concept: PBE

Equilibrium analysis

Full information benchmark

Suppose there is **no uncertainty** – players know each other's types from the start.
Then, equilibrium outcomes would be:

- ▶ Hard type v. hard type $\rightarrow SQ$
- ▶ Hard type v. soft type $\rightarrow$ preferred alternative of hard type
- ▶ Soft type v. soft type $\rightarrow$ preferred alternative of whomever gets the first commitment opportunity

Equilibrium strategies with uncertainty

- ▶ **Result 1:** Hard types commit ASAP to their preferred policy
- ▶ **Result 2:** Soft types *who know their opponent's type* also act ASAP
- ▶ *Uninformed* soft types want to delay long enough that hard types screen out, but not so long that soft types realize they're stalling
 - ▶ **Preemption** (better chance of getting favorite policy, but could alienate an obstinate opponent) vs **caution** (more time to learn, but risk getting beaten to the punch)
 - ▶ **Result 3:** The time at which they become willing to play hardball balances these two incentives

Equilibrium strategies with uncertainty

- ▶ A **symmetric setting** of the model helps us isolate and understand this result:
 - ▶ Identical priors: $p_a = p_b$
 - ▶ Identical issue salience: $u_a^s(A) = u_b^s(B)$, $u_a^s(B) = u_b^s(A)$, and $u_a^s(SQ) = u_b^s(SQ)$
 - ▶ Identical within-type arrival rates: $\mu_s^a = \mu_s^b$, $\mu_h^a = \mu_h^b$

Proposition

Consider a history where neither group has committed to an alternative.

In the unique PBE, a soft group commits to its preferred alternative iff $t > T^$.*

Outcomes

What are the **patterns of outcomes** implied by equilibrium strategies?

- ▶ Both groups are hard types $\implies$ compromise is impossible; status quo remains in place
- ▶ Both groups are soft types $\implies$ first group to receive an opportunity **after** T^* gets its preferred policy
- ▶ One group is a hard type and the other is a soft type $\implies$ hard types *usually* get their preferred policy
 - ▶ **But not always!** Sometimes soft types preempt hard types, resulting in the status quo although a compromise was possible

What's driving the result?

$$T_a^* = -\frac{1}{\mu_h} \ln \left(\frac{1-p}{p} \underbrace{\frac{1}{2} \frac{\mu_h + \mu_s}{\mu_h} \frac{u_a(A) - u_a(B)}{u_a(B) - u_a(SQ)}}_{\text{Threshold posterior belief}} \right)$$

- ▶ **Broad level:** Equilibrium delay is pinned down by a **threshold posterior belief**
- ▶ **Specific level:** Factors serve as vehicles for beliefs, determining the **value of learning more** or the **speed of learning**

Comparative statics on T^* : Beliefs, relative desirability

$$T_a^* = -\frac{1}{\mu_h} \ln \left(\underbrace{\frac{1-p}{p}}_{\text{Relative priors}} \frac{1}{2} \frac{\mu_h + \mu_s}{\mu_h} \frac{u_a(A) - u_a(B)}{u_a(B) - u_a(SQ)} \right)$$

$T^* \rightarrow 0$ when the soft group has...

- ▶ **Strong priors** that the opponent is a soft rather than a hard type ($\frac{1-p}{p}$ is high)
- ▶ *Intuition:* Beliefs have less distance to travel to the threshold

Comparative statics on T^* : Beliefs, relative desirability

$$T_a^* = -\frac{1}{\mu_h} \ln \left(\frac{1-p}{p} \frac{1}{2} \frac{\mu_h + \mu_s}{\mu_h} \underbrace{\frac{u_a(A) - u_a(B)}{u_a(B) - u_a(SQ)}}_{\text{Relative desirability}} \right)$$

$T^* \rightarrow 0$ when the soft group has...

- ▶ **Strong relative preference** for their favorite outcome $\left(\frac{u(A) - u(B)}{u(B) - u(SQ)} \right)$ is high
- ▶ *Intuition:* less benefit to learning more about the opponent

Comparative statics on T^* : Political frictions

$$T_a^* = -\frac{1}{\mu_h} \ln \left(\frac{1-p}{p} \frac{1}{2} \frac{\mu_h + \mu_s}{\mu_h} \frac{u_a(A) - u_a(B)}{u_a(B) - u_a(SQ)} \right)$$

$T^* \rightarrow 0$ when there are...

- ▶ Faster **arrivals of hardball opportunities for soft types**
- ▶ *Intuition:* Beliefs threshold becomes easier to reach

Comparative statics on T^* : Political frictions

$$T_a^* = -\frac{1}{\mu_h} \ln \left(\frac{1-p}{p} \frac{1}{2} \frac{\mu_h + \mu_s}{\mu_h} \frac{u_a(A) - u_a(B)}{u_a(B) - u_a(SQ)} \right)$$

$T^* \rightarrow 0$ when there are...

- ▶ Faster **arrivals of hardball opportunities** for hard types
- ▶ *Intuition:* Easier to distinguish strategic from accidental stalling

What factors drive learning and delay?

- ▶ Each side's type (*willingness to compromise*)
- ▶ Soft types' **knowledge and beliefs** about their opponents (*reputation*)
- ▶ Soft types' **relative preference** for their favorite outcome (*issue salience*)
- ▶ **Arrival rate of hardball opportunities** (*political frictions*)

How do these factors affect strategic behavior when they differ *between* players?

One-sided incomplete information

One-sided incomplete information

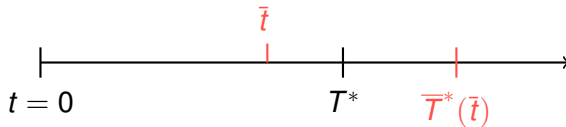
- ▶ At $\bar{t}$, one player is revealed as a soft type
- ▶ **Direct effect:** Opponent fully informed, therefore wants to commit irrespective of opponent's own type
- ▶ **Indirect effect:** The revealed player has **less information** from which to learn about the opponent
 - ▶ Posterior beliefs converge more slowly to the threshold $\implies$ revealed player delays longer

Equilibrium strategies with asymmetric information

Proposition

Suppose a soft group's type is revealed at time $\bar{t} < T^$.*

Then, there exists a unique $\bar{T}^(\bar{t}) > T^*$ such that a soft group will commit to its preferred alternative iff $t > \bar{T}^*(\bar{t})$.*



We can apply this intuition to other asymmetries in the model!

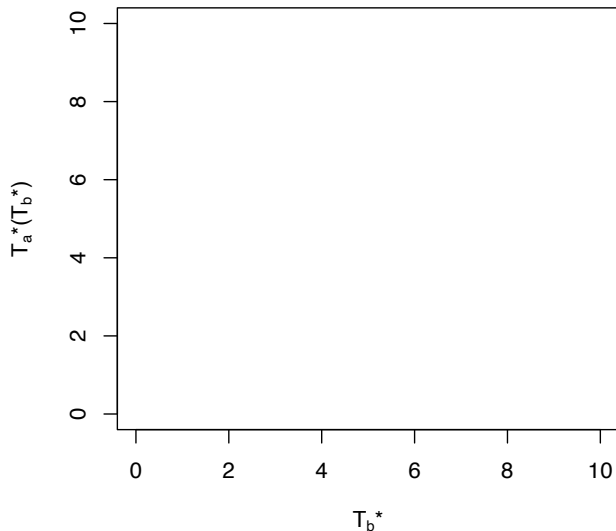
Asymmetric setting

Asymmetric priors

- ▶ Suppose players' priors are asymmetric: player a has a **higher reputation for being unwilling to compromise**
- ▶ Consequence: $T_a^* < T_b^*$
- ▶ Direct effect: After T_a^* , player b knows a is trying to commit
 - ▶ **Indirect effect 1:** This slows down learning for player b , causing player b to delay longer
 - ▶ **Indirect effect 2:** This affords player a some “slack,” since exercising some more caution carries clear benefits and very little risk

Visualizing with best responses: Symmetric case

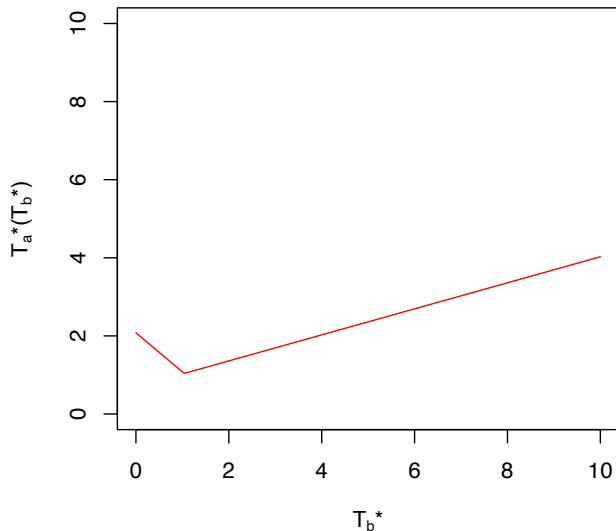
- ▶ Consider a soft type of group a 's best response to how much a soft type of group b delays
- ▶ In other words, $T_a^*(T_b)$



▶ Skip to welfare

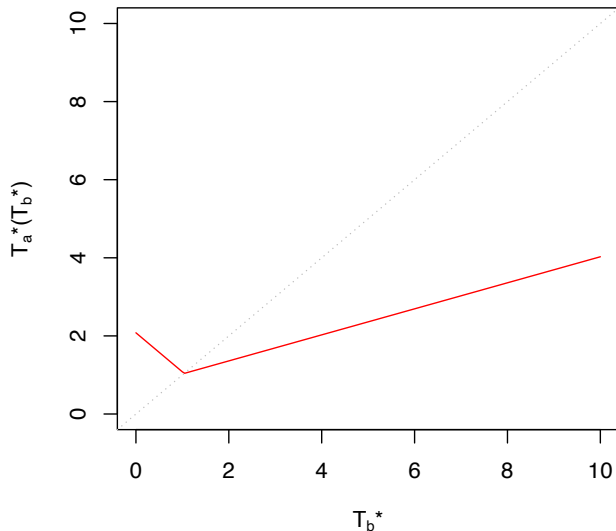
Visualizing with best responses: Symmetric case

- ▶ Group a 's best response,
 $T_a^*(T_b^*)$



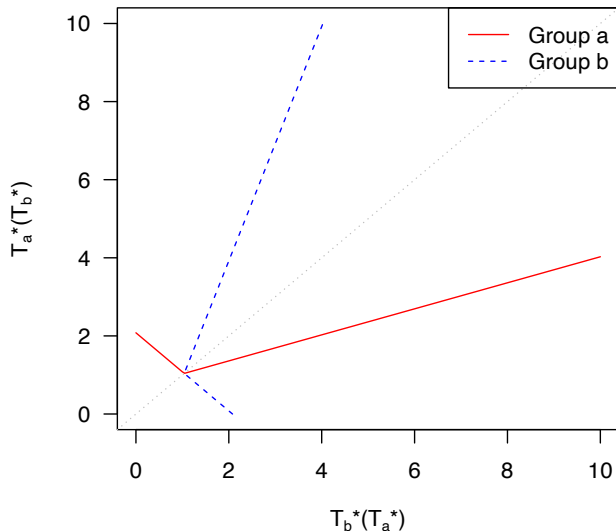
Visualizing with best responses: Symmetric case

- ▶ Group a 's best response, $T_a^*(T_b^*)$
- ▶ Best response kinks at the 45-degree line



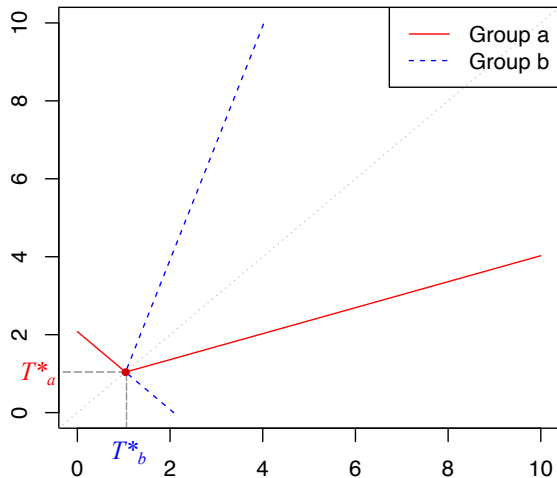
Visualizing with best responses: Symmetric case

- ▶ Group a 's best response, $T_a^*(T_b^*)$
- ▶ Group b 's best response, $T_b^*(T_a^*)$



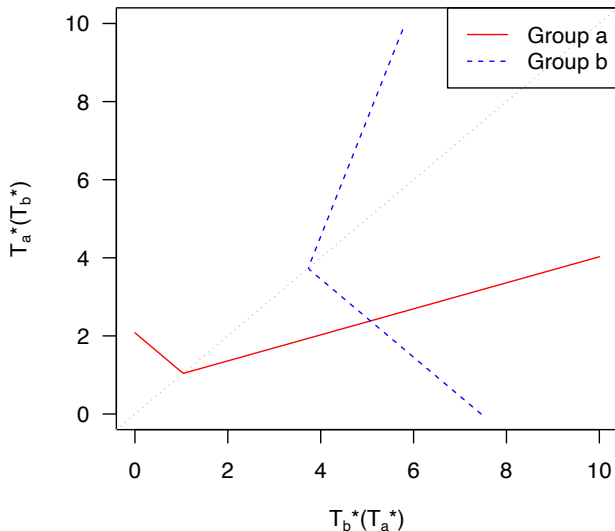
Visualizing with best responses: Symmetric case

- ▶ Group a 's best response, $T_a^*(T_b^*)$
- ▶ Group b 's best response, $T_b^*(T_a^*)$
- ▶ Minimum delay in eqm by both players



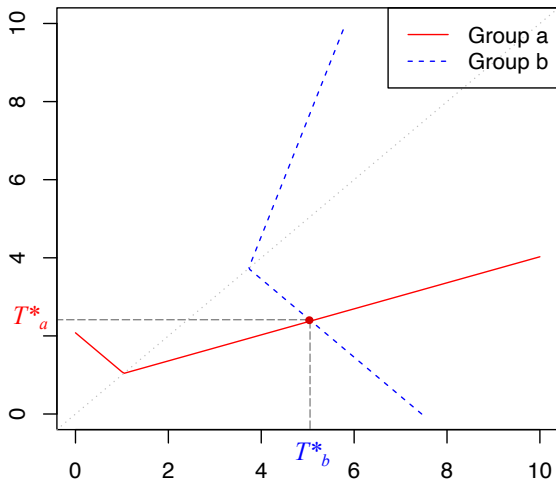
Visualizing with best responses: Asymmetric setting

- ▶ Suppose player a 's has a **higher reputation for being a hard type** ($p_b > p_a$)
- ▶ Consequence: player b 's best response shifts



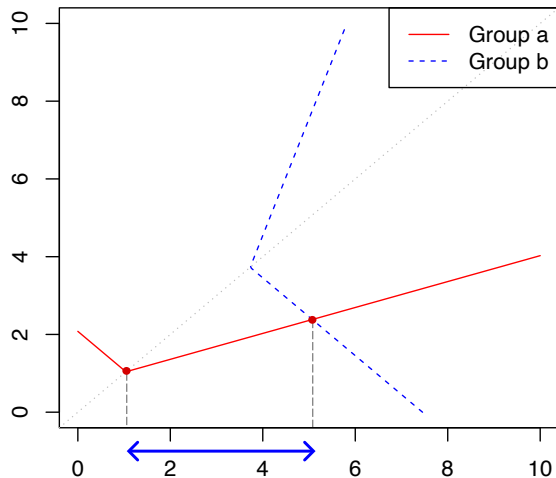
Visualizing with best responses: Asymmetric setting

- ▶ Consequence: $T_b^* > T_a^*$
- ▶ More likely that player a gets their preferred policy



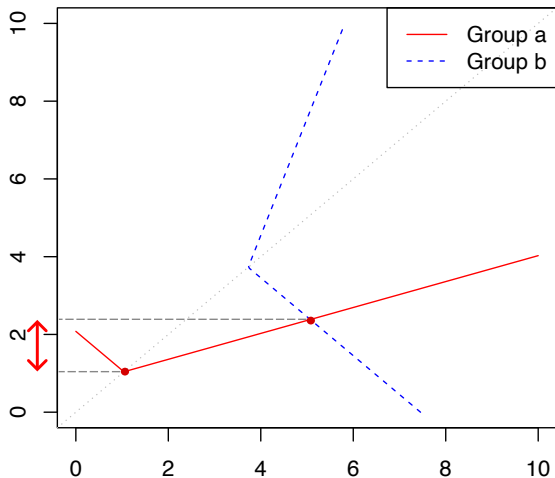
Visualizing with best responses: Indirect effect 1

- **Indirect effect 1:** Player *b* delays longer because of slower learning



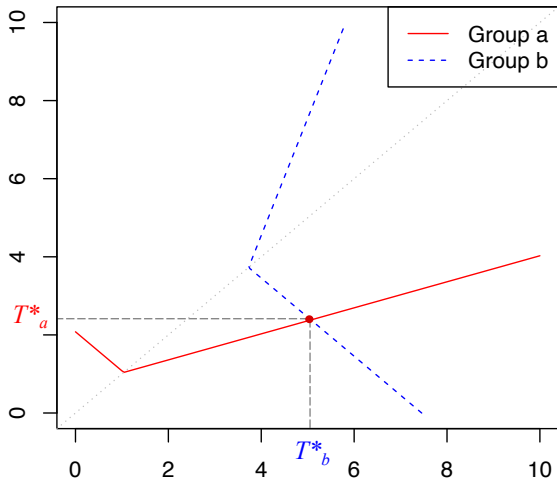
Visualizing with best responses: Indirect effect 2

- **Indirect effect 2:** Player *a* also delays longer than in the symmetric case because of “slack”



A comment on uniqueness

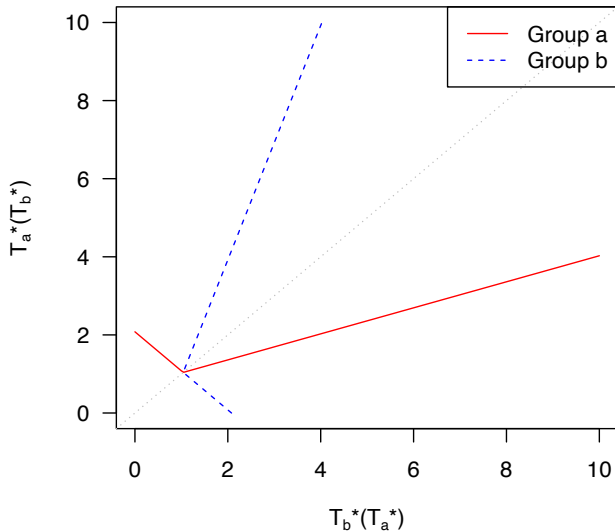
- ▶ Because of both indirect effects, best responses have a check shape
- ▶ This shape guarantees equilibrium uniqueness in the space of threshold strategies
- ▶ With more work you can rule out PBEs in non-threshold strategies
→ generic uniqueness



A comment on comparative statics

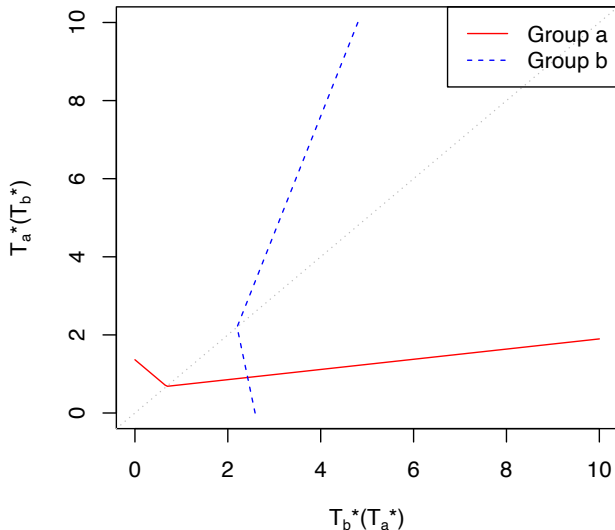
- ▶ Priors and relative desirability only shift player a 's best response
- ▶ However, μ_s^a factors into both players' best responses

▶▶ Best response expressions



A comment on comparative statics

- ▶ Decrease μ_s^a : player a get opportunities less frequently on average
- ▶ Player b delays longer because it has a “technological” advantage
- ▶ Meanwhile, player a delays less to make up for its slower arrival rate



A brief comment regarding welfare

- ▶ How do equilibrium strategies generate **avoidable welfare loss**?
 - ▶ “**Preemption** (better chance of getting favorite policy, but **could alienate an obstinate opponent**) vs **caution** (more time to learn, but risk getting beaten to the punch)”
 - ▶ Mistakes of preemption are the *inevitable consequence* of this trade-off
- ▶ In the symmetric setting, equilibrium behavior affects welfare **solely** through this channel → comparative statics which extend delay improve welfare
- ▶ Asymmetry extends delay for both players, reducing avoidable miscoordination

Extension: Introducing a leaky environment

- ▶ Previously, we *imposed* that one player's type becomes publicly known at $\bar{t}$
- ▶ What if instead “leaks” might occur at any point in the game? (E.g. gossip, press investigations, reconnaissance)
- ▶ Suppose leaks occur according to a Poisson process, essentially inducing a distribution over $\bar{t}$
- ▶ Do players delay more or less? What are the consequences for successful coordination?
- ▶ Depends on if leaks are correlated with type – more leaks of “hard types” facilitate screening and reduce delay; whereas leaks of soft types impede screening and extend delay

Conclusion

Recap of key ideas

- ▶ **How do groups coordinate in the presence of internal disagreements?**
- ▶ *Starting point:* Uncertainty over willingness to compromise, and political frictions generating opportunities to “hardball”
- ▶ *Mechanism:* **Endogenous learning** about willingness to compromise on the basis of behavior
- ▶ *Enriching the story:* Asymmetries create additional cautionary incentives for at least one player

Connections to outcomes and welfare

- ▶ How did learning and strategic behavior translate into **negotiation outcomes** and **players' welfare**?
- ▶ Longer delay implies lower chance of getting favorite policy, but also lower chance of *avoidable* miscoordination
- ▶ Factors like reputation, issue salience, and political frictions can change the pace or value of learning, changing the likelihood of avoidable miscoordination

Why does this matter in politics?

- ▶ Gives us a way of thinking about how united fronts emerge, or fail to emerge, from times of uncertainty and inaction
- ▶ Systematizes intuitions about guessing your opponent's intentions
- ▶ Elucidates the dynamic incentives inherent in the coordination process
- ▶ Pertains to fundamental questions about cooperation and conflict: Alliances in legislatures, rebel groups, authoritarian elites, and international organizations

Thanks!

eyao@princeton.edu

Supplementary material

- ▶ Leaky environments ▶ [Link](#)
- ▶ Eqm expressions in leaky environments ▶ [Link](#)
- ▶ Welfare ▶ [Link](#)

Leaky informational environments

Leaky informational environments

- ▶ Suppose that a player's type is publicly revealed $\sim \text{Poisson}(\lambda)$, where λ can vary by player and type
- ▶ This induces a **distribution over $\bar{t}$**
- ▶ Changes learning, and introduces **new incentives to hedge against leaks**

▶ [Return to main slides](#)

Comparative statics on leaks

$$T_a^* = \frac{1}{\lambda_s - (\lambda_h + \mu_h)} \left[\ln \left(\frac{1}{2} \cdot \frac{1-p}{p} \cdot \frac{(\lambda_h + \mu_h) + \mu_s}{(\lambda_h + \mu_h)} \frac{u_a(A) - u_a(B)}{u_a(B) - u_a(SQ)} \right) \right]$$

- ▶ Delay is increasing in λ_s – learning is slower
- ▶ Delay is increasing in λ_h – screening is faster

Leaks extend delay and reduce avoidable miscoordination

- ▶ Both T^* and $\overline{T}^*$ are increasing in λ_s

Leaks extend delay and reduce avoidable miscoordination

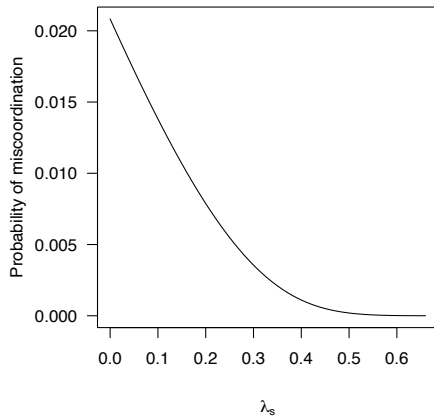
- ▶ Both T^* and $\overline{T}^*$ are increasing in λ_s
- ▶ Higher λ_s makes it more likely that you will be leaked, landing in the one-sided information case

Leaks extend delay and reduce avoidable miscoordination

- ▶ Both T^* and $\overline{T}^*$ are increasing in λ_s
- ▶ Higher λ_s makes it more likely that you will be leaked, landing in the one-sided information case
- ▶ All of these effects drive at increased delay, and fewer mistakes of preemption

Leaks extend delay and reduce avoidable miscoordination

- ▶ Both T^* and $\overline{T}^*$ are increasing in λ_s
- ▶ Higher λ_s makes it more likely that you will be leaked, landing in the one-sided information case
- ▶ All of these effects drive at increased delay, and fewer mistakes of preemption
- ▶ Result wouldn't be obvious without the model! λ_s is a “cost,” but increasing it makes soft types globally better off



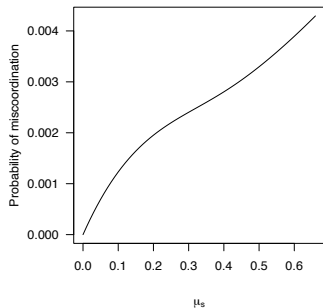
(a) Miscoordination is less likely when soft types are leaked more frequently

Leaks condition the effects of political frictions

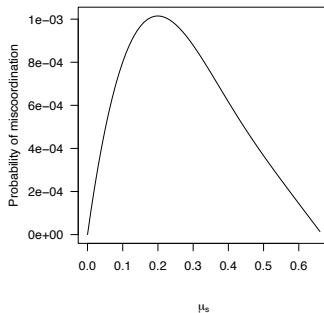
- ▶ When neither group has been leaked, higher μ_s means T^* is shorter
- ▶ **In the one-sided information case**, this result breaks down – direction of the comparative static depends on the time of leak
- ▶ The total effect on avoidable miscoordination is based on **how likely it is that the leak scenario occurs**

▶ Return to main slides

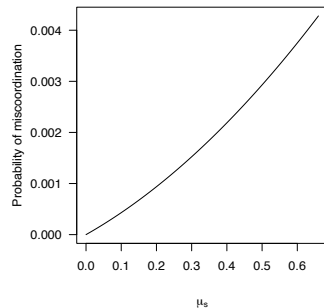
Disaggregating the effects of μ_s on miscoordination



(b) Aggregate effect of μ_s on miscoordination



(c) One group has been leaked



(d) Neither group has been leaked

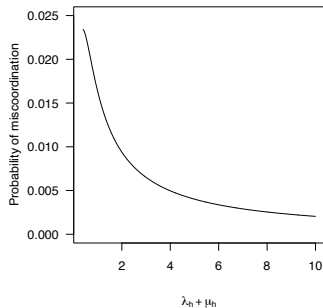
► [Return to main slides](#)

Leaks condition the effect of $\lambda_h + \mu_h$

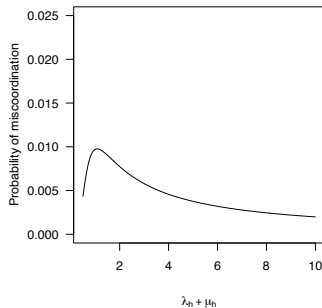
- ▶ As $\lambda_h + \mu_h \rightarrow \infty$, hard types screen out almost instantaneously, so soft types never make a mistake
- ▶ **If one group is leaked**, this result breaks down at lower values of $\lambda_h + \mu_h$
- ▶ Soft types exhibit **rational impatience**: Delay is expensive, so you accept a higher risk of making a mistake of preemption

▶ Return to main slides

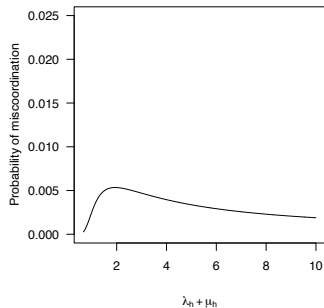
$\lambda_h + \mu_h$ reduces miscoordination... eventually



(e) $\lambda_s = \frac{1}{16}$



(f) $\lambda_s = \frac{1}{6}$



(g) $\lambda_s = \frac{1}{3}$

► [Return to main slides](#)

Takeaways from leaky environments

- ▶ Leaks are good on their own: They reduce mistakes of preemption and increase total welfare
- ▶ But when leaks are frequent, the effects of other factors become less predictable: μ_s and μ_h have non-monotonic effects on welfare
- ▶ Key: Leaks create a threat of **one-sided asymmetric information**
- ▶ When this threat is high, players have different (more preemptive) incentives

▶ [Return to main slides](#)

- ▶ Introducing the possibility of leaks complicate the incentives that come from changing other factors, such as the speed of commitments
- ▶ When leaks are likely, they force players to hedge against the probability that they will be leaked
- ▶ These can create new preemptive incentives, making it more difficult to ascertain how changes in the speed of commitment opportunities will affect welfare

▶ [Return to main slides](#)

Supplementary equilibrium expressions

Characterizing threshold beliefs with leaks

$$\frac{\mathbb{P}(b \text{ is a hard type})}{\mathbb{P}(b \text{ is a soft type})} = \frac{1}{2} \frac{(\lambda_h + \mu_h) + \mu_s}{(\lambda_h + \mu_h)} \frac{u_a(A) - u_a(B)}{u_a(B) - u_a(SQ)}$$

► Return to main slides

Characterizing T^* with leaks

$$T^* = \frac{1}{\lambda_s - (\lambda_h + \mu_h)} \left[\ln \left(\underbrace{\frac{1}{2} \cdot \frac{1-p}{p} \cdot \frac{(\lambda_h + \mu_h) + \mu_s}{(\lambda_h + \mu_h)} \frac{u_a(A) - u_a(B)}{u_a(B) - u_a(SQ)}}_{\text{Threshold beliefs}} \right) \right]$$

► Return to main slides

Welfare: Symmetric case

Welfare

- ▶ How does equilibrium behavior and comparative statics map onto welfare?
- ▶ Object of analysis: (soft) player's ex ante *expected* utility:

$$(1 - p) \left[\frac{u_a(A) + u_a(B)}{2} \right] + (p) \left[\mathbb{P}(\text{avoidable miscoordination}) u_a(SQ) + \left(1 - \mathbb{P}(\text{avoidable miscoordination}) \right) u_a(B) \right]$$

▶ [Return to main slides](#)

Inefficiencies come from avoidable miscoordination

- ▶ **Inefficient and avoidable miscoordination** occurs when a soft type commits to its preferred alternative, but the opponent is a hard type
- ▶ Avoidable miscoordination is decreasing in delay
- ▶ Factors that unilaterally decrease delay without affecting other components of welfare are welfare-increasing
- ▶ The effect of **political frictions and leaks** (λ, μ) works exclusively through the channel of avoidable miscoordination

▶ [Return to main slides](#)

Avoidable miscoordination

- ▶ Given some arbitrary amount of delay T , the probability of avoidable miscoordination is

$$e^{(-\mu_h - \lambda_h)T} \frac{\mu_s}{\mu_s + \mu_h + \lambda_h}$$

which is decreasing in T .

Avoidable miscoordination

- Plugging in the equilibrium expressions for T^* , the probability of avoidable miscoordination is

$$\begin{aligned} & \frac{\mu_s}{\lambda_h + \mu_h + \mu_s} \left[\left(\frac{1-p}{2p} \frac{u(A) - u(B)}{u(B) - u(SQ)} \frac{\lambda_h + \mu_h + \mu_s}{\lambda_h + \mu_h} \right)^{\frac{\lambda_s + \lambda_h + \mu_h}{\lambda_h + \mu_h - \lambda_s}} \right. \\ & + \frac{\lambda_s(\lambda_h + \mu_h - \lambda_s - \mu_s)}{\mu_s(\mu_h + \lambda_h) - \lambda_s(\lambda_h + \mu_h - \lambda_s - \mu_s)} \left(\left(\frac{1-p}{2p} \frac{u(A) - u(B)}{u(B) - u(SQ)} \frac{\lambda_h + \mu_h + \mu_s}{\lambda_h + \mu_h} \right)^{\frac{\lambda_s + \lambda_h + \mu_h}{\lambda_h + \mu_h - \lambda_s}} \right. \\ & \left. \left. - \left(\frac{1-p}{2p} \frac{u(A) - u(B)}{u(B) - u(SQ)} \frac{\lambda_h + \mu_h + \mu_s}{\lambda_h + \mu_h} \right)^{\frac{\lambda_h + \mu_h}{\lambda_h + \mu_h - \lambda_s - \mu_s}} \right) \right] \end{aligned}$$

Asymmetric setting

Best responses

$$T_i^*(T_j^*)|T_j^* < K_i = \frac{1}{\lambda_s^j + \mu_s^j - \Lambda^j} \left(\ln \left[\frac{u_i(B) - u_i(A)}{u_i(A) - u_i(SQ)} \frac{\mu_s^j}{\mu_s^j + \mu_s^i} \frac{1 - p_i}{p_i} \frac{\Lambda^j + \mu_s^i}{\Lambda^j} \right] + \mu_s^j T_j^* \right)$$

$$T_i^*(T_j^*)|T_j^* > K_i = \frac{1}{\lambda_s^j - \Lambda^j - \mu_s^i} \left(\ln \left[\frac{u_i(B) - u_i(A)}{u_i(A) - u_i(SQ)} \frac{\mu_s^j}{\mu_s^j + \mu_s^i} \frac{1 - p_i}{p_i} \frac{\Lambda^j + \mu_s^i}{\Lambda^j} \right] - \mu_s^i T_j^* \right)$$

where

$$K_i := \frac{\mu_s^i + \mu_s^j}{\mu_s^i(\lambda_s^i - \Lambda^i - \mu_s^j) + \mu_s^j(\lambda_s^i + \mu_s^i - \Lambda^i)} \ln \left[\frac{u_j(A) - u_j(B)}{u_j(B) - u_j(SQ)} \frac{\mu_s^i}{\mu_s^i + \mu_s^j} \frac{1 - p_j}{p_j} \frac{\Lambda^i + \mu_s^j}{\Lambda^i} \right]$$

$$\Lambda^i = \lambda_h^i + \mu_h^i$$

► [Return to main slides](#)