

Non-monotonic Disclosure in Policy Advice

Anna Denisenko (U. Chicago), Catherine Hafer (NYU), and Dimitri Landa (NYU)

17 Nov 2025

Motivation

Strategic communications between policymakers and bureaucratic agencies

Motivation

Strategic communications between policymakers and bureaucratic agencies

- Communications often occur with **verifiable information**
 - internal norms or rules

Motivation

Strategic communications between policymakers and bureaucratic agencies

- Communications often occur with **verifiable information**
 - internal norms or rules
- Policymakers (elected officials) and bureaucrats preferences are frequently misaligned
 - bureaucrats less affected by short-term public opinion volatility
 - bureaucrats like status quo

Motivation

Strategic communications between policymakers and bureaucratic agencies

- Communications often occur with **verifiable information**
 - internal norms or rules
- Policymakers (elected officials) and bureaucrats preferences are frequently misaligned
 - bureaucrats less affected by short-term public opinion volatility
 - bureaucrats like status quo

Disclosure Games

- Preference misalignment under verifiable information → full disclosure (Milgrom (1981), Grossman (1981))

Motivation

Strategic communications between policymakers and bureaucratic agencies

- Communications often occur with **verifiable information**
 - internal norms or rules
- Policymakers (elected officials) and bureaucrats preferences are frequently misaligned
 - bureaucrats less affected by short-term public opinion volatility
 - bureaucrats like status quo

Disclosure Games

- Preference misalignment under verifiable information → full disclosure (Milgrom (1981), Grossman (1981))
 - monotonicity of Sender preferences in Receiver action

Some Results

- ① When ex-ante preferences of sender and receiver sufficiently co-align, disclosure may be partial, contrary to standard *unraveling* results
 - no uncertainty about informedness of Sender

Some Results

- ① When ex-ante preferences of sender and receiver sufficiently co-align, disclosure may be partial, contrary to standard *unraveling* results
 - no uncertainty about informedness of Sender
- ② For sufficient co-alignment, there is multiplicity of Sequential Equilibria
 - introduce a novel equilibrium refinement: belief stability
 - belief-stable equilibria exhibit different comparative statics than belief-unstable equilibria

Some Results

- ① When ex-ante preferences of sender and receiver sufficiently co-align, disclosure may be partial, contrary to standard *unraveling* results
 - no uncertainty about informedness of Sender
- ② For sufficient co-alignment, there is multiplicity of Sequential Equilibria
 - introduce a novel equilibrium refinement: belief stability
 - belief-stable equilibria exhibit different comparative statics than belief-unstable equilibria
- ③ Robust prediction that more ex-ante preference divergence yields more information disclosure, contrary to cheap-talk results

Our Contributions

- Full disclosure in games of verifiable advice:
 - **Milgrom (1981), Grossman (1981), Milgrom (2008)**
 - **Seidmann and Winter (1997)**
 - objective function concave in action
 - sender's utility more state-dependent than receiver's
- Partial disclosure in games of verifiable advice
 - uninformed sender **Dye (1985), Jung and Kwon (1988)**
 - uncertainty about S's preferences **Wolinsky (2003), Dziuda (2011)**
 - multidimensional advice **Callander, Lambert and Matouschek (2021)**
 - disclosure reward **Denisenko, Hafer and Landa (2024)**
- Games of communication within hierarchy (cheap talk)
 - divergence in preferences → worse communication: seminal paper by **Crawford and Sobel (1982), Gilligan and Kreihbiel (1987), Austen-Smith (1990, 1993)**
 - **Callander (2008)**

Road Map

- ① Introduction
- ② Uniform Prior, Quadratic Preferences, State-independent Sender Preferences
 - Game Structure
 - Equilibrium Characterization
 - Effects of Preference Divergence
 - Belief Stability Equilibrium Refinement
- ③ General Model
 - Equilibria
 - Comparative Statics
- ④ Robustness
- ⑤ Summary

Actors and Timing

Two players: Agency (it) and Policymaker (she).

Actors and Timing

Two players: Agency (it) and Policymaker (she).

①	Nature determines realization of the state of the world (ω)	$\omega \sim U[-1, 1]$
---	---	------------------------



Actors and Timing

Two players: Agency (it) and Policymaker (she).

①	Nature determines realization of the state of the world (ω)	$\omega \sim U[-1, 1]$
②	Agency observes state (ω)	ω



Actors and Timing

Two players: Agency (it) and Policymaker (she).

①	Nature determines realization of the state of the world (ω)	$\omega \sim U[-1, 1]$
②	Agency observes state (ω)	ω
③	Agency chooses message (m) to send to Policymaker	$m \in \{\omega, \emptyset\}$



Actors and Timing

Two players: Agency (it) and Policymaker (she).

①	Nature determines realization of the state of the world (ω)	$\omega \sim U[-1, 1]$
②	Agency observes state (ω)	ω
③	Agency chooses message (m) to send to Policymaker	$m \in \{\omega, \emptyset\}$
④	Policymaker observes m and chooses policy (p)	$p \in \mathbb{R}$



Uniform Prior and Quadratic Preferences: Payoffs and Solution Concept

- Agency:

$$u_A(p) = -(p - i)^2$$

Uniform Prior and Quadratic Preferences: Payoffs and Solution Concept

- Agency:

$$u_A(p) = -(p - i)^2,$$

where i is Agency's ideal point.

Uniform Prior and Quadratic Preferences: Payoffs and Solution Concept

- Agency:

$$u_A(p) = -(p - i)^2,$$

where i is Agency's ideal point.

- Policymaker:

$$u_P(p) = -(p - \omega)^2,$$

Uniform Prior and Quadratic Preferences: Payoffs and Solution Concept

- Agency:

$$u_A(p) = -(p - i)^2,$$

where i is Agency's ideal point.

- Policymaker:

$$u_P(p) = -(p - \omega)^2,$$

$\Rightarrow |i|$ is **ex-ante** actors' preference divergence.

Uniform Prior and Quadratic Preferences: Payoffs and Solution Concept

- Agency:

$$u_A(p) = -(p - i)^2,$$

where i is Agency's ideal point.

- Policymaker:

$$u_P(p) = -(p - \omega)^2,$$

$\Rightarrow |i|$ is **ex-ante** actors' preference divergence.

Solution Concept: Sequential Equilibrium.

Road Map

- ① Introduction
- ② Uniform Prior, Quadratic Preferences, State-independent Sender Preferences
 - Game Structure
 - **Equilibrium Characterization**
 - Effects of Preference Divergence
 - Belief Stability Equilibrium Refinement
- ③ General Model
 - Equilibria
 - Comparative Statics
- ④ Robustness
- ⑤ Summary

Equilibrium Characterization

In every equilibrium

Policymaker

- $p^*(m = \omega) = \omega$ when $m \neq \emptyset$;
- $p^*(m = \emptyset) = x^* \equiv E[\omega | m^*(\omega) = \emptyset]$,
where $m^*(\omega)$ is A's eq. disclosure strategy.

Equilibrium Characterization

In every equilibrium

Policymaker

- $p^*(m = \omega) = \omega$ when $m \neq \emptyset$;
- $p^*(m = \emptyset) = x^* \equiv E[\omega | m^*(\omega) = \emptyset]$,
where $m^*(\omega)$ is A's eq. disclosure strategy.

Agency

- discloses ω when
 $\omega \in [i - \sqrt{(x^* - i)^2}, i + \sqrt{(x^* - i)^2}] \cap [-1, 1]$;
- conceals ω otherwise.

Equilibrium Characterization

In every equilibrium

Policymaker

- $p^*(m = \omega) = \omega$ when $m \neq \emptyset$;
- $p^*(m = \emptyset) = x^* \equiv E[\omega | m^*(\omega) = \emptyset]$,
where $m^*(\omega)$ is A's eq. disclosure strategy.

Agency

- discloses ω when
 $\omega \in [i - \sqrt{(x^* - i)^2}, i + \sqrt{(x^* - i)^2}] \cap [-1, 1]$;
- conceals ω otherwise.

$i \geq 0 \rightarrow$ disclose $\omega \in [x^*, 2 \cdot i - x^*] \cap [-1, 1]$;

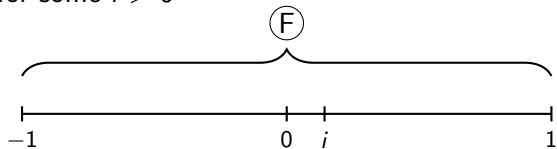
$i \leq 0 \rightarrow$ disclose $\omega \in [2 \cdot i - x^*, x^*] \cap [-1, 1]$.

Equilibrium Disclosure Strategies

There can be a *maximum* of three disclosure strategies supported in equilibrium

① Full disclosure (F)

Disclosure intervals for some $i > 0$



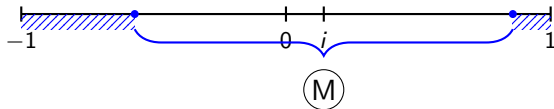
Hatched areas – no disclosure

Equilibrium Disclosure Strategies

There can be a *maximum* of three disclosure strategies supported in equilibrium

- ① Full disclosure (F)
- ② **Partial disclosure:**
 - **More Expansive disclosure strategy (M)**

Disclosure intervals for some $i > 0$:



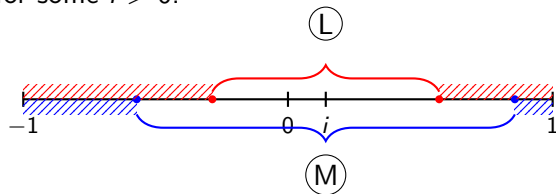
Hatched areas – no disclosure

Equilibrium Disclosure Strategies

There can be a *maximum* of three disclosure strategies supported in equilibrium

- ① Full disclosure (F)
- ② **Partial disclosure:**
 - **More Expansive disclosure strategy (M),**
 - **Less Expansive disclosure strategy (L).**

Disclosure intervals for some $i > 0$:



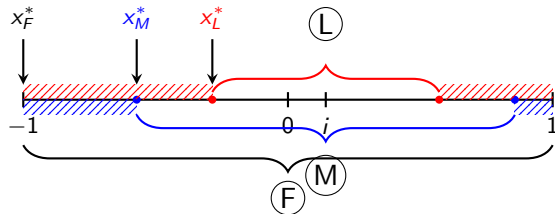
Hatched areas – no disclosure

Equilibria

There can be a *maximum* of three equilibria

- ① Full disclosure equilibrium;
- ② Partial disclosure equilibria:
 - More Expansive equilibrium,
 - Less Expansive equilibrium.

Disclosure intervals for some $i > 0$:



Hatched areas – no disclosure

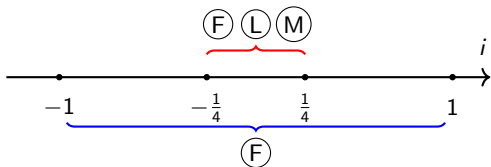
Road Map

- ① Introduction
- ② Uniform Prior, Quadratic Preferences, State-independent Sender Preferences
 - Game Structure
 - Equilibrium Characterization
 - Effects of Preference Divergence on
 - Existence of Partial Disclosure Equilibrium
 - Disclosure
 - Belief Stability Equilibrium Refinement
- ③ General Model
 - Equilibria
 - Comparative Statics
- ④ Robustness
- ⑤ Summary

Effect of Preference Divergence ($|i|$) on Full Disclosure Equilibrium Uniqueness

Prop.1

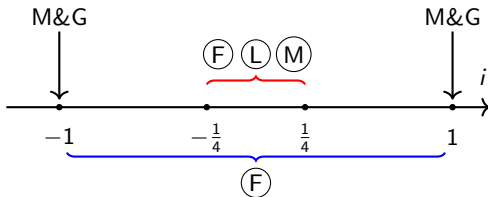
- ① For all i there exists full disclosure equilibrium;
- ② If and only if $|i| \leq 1/4$, there are two partial disclosure equilibria: **less expansive** and **more expansive**.



Effect of Preference Divergence ($|i|$) on Full Disclosure Equilibrium Uniqueness

Prop.1

- ① For all i there exists full disclosure equilibrium;
- ② If and only if $|i| \leq 1/4$, there are two partial disclosure equilibria: **less expansive** and **more expansive**.



Effect of Preference Divergence ($|i|$) on Equilibrium Disclosure

Assume $i \geq 0 \rightarrow$

Agency discloses ω to PM when

$$\omega \in [x^*, 2 \cdot i - x^*] \cap [-1, 1],$$

and conceals information otherwise.

Effect of Preference Divergence ($|i|$) on Equilibrium Disclosure

Assume $i \geq 0 \rightarrow$

Agency discloses ω to PM when

$$\omega \in [x^*, 2 \cdot i - x^*] \cap [-1, 1],$$

and conceals information otherwise.

Preference divergence ($|i|$) has **direct**

Effect of Preference Divergence ($|i|$) on Equilibrium Disclosure

Assume $i \geq 0 \rightarrow$

Agency discloses ω to PM when

$$\omega \in [x^*, 2 \cdot i - x^*] \cap [-1, 1],$$

and conceals information otherwise.

Preference divergence ($|i|$) has **direct** and **indirect** effects on disclosure.

Effect of Preference Divergence ($|i|$) on Equilibrium Disclosure

Assume $i \geq 0 \rightarrow$

Agency discloses ω to PM when

$$\omega \in [x^*, 2 \cdot i - x^*] \cap [-1, 1],$$

and conceals information otherwise.

Preference divergence ($|i|$) has **direct** and **indirect** effects on disclosure.

- **Direct** effect always (*weakly*) improves communication between A and PM
- **Indirect** effect
 - $\rightarrow$ Improves communication in **less expansive** equilibrium

Effect of Preference Divergence ($|i|$) on Equilibrium Disclosure

Assume $i \geq 0 \rightarrow$

Agency discloses ω to PM when

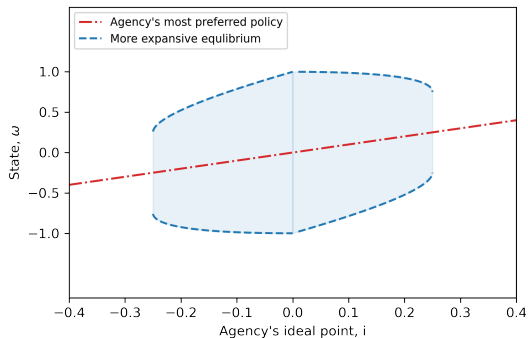
$$\omega \in [x^*, 2 \cdot i - x^*] \cap [-1, 1],$$

and conceals information otherwise.

Preference divergence ($|i|$) has **direct** and **indirect** effects on disclosure.

- **Direct** effect always (*weakly*) improves communication between A and PM
- **Indirect** effect
 - $\rightarrow$ Improves communication in **less expansive** equilibrium
 - $\rightarrow$ Reduces communication in **more expansive** equilibrium

Effect of Preference Divergence ($|i|$) in More Expansive Eq'm

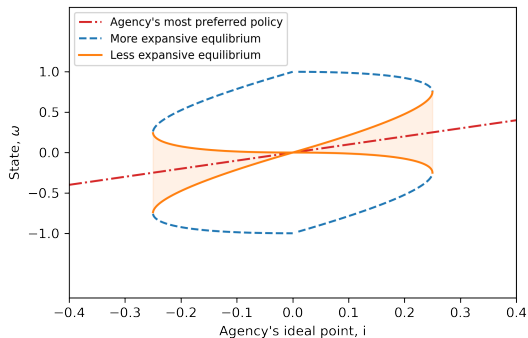


Prop.2

Communication between actors

→ *deteriorates* in $|i|$ in more expansive equilibrium;

Effect of Preference Divergence ($|i|$) in Less Expansive Eq'm

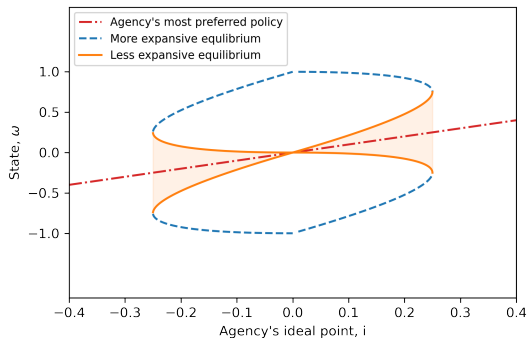


Prop.2

Communication between actors

- *deteriorates* in $|i|$ in more expansive equilibrium;
- *improves* in $|i|$ in less expansive equilibrium;

Effect of Preference Divergence ($|i|$) in Less Expansive Eq'm



Comparative Statics Underlying Intuition

Prop.2

Communication between actors

- *deteriorates* in $|i|$ in more expansive equilibrium;
- *improves* in $|i|$ in less expansive equilibrium; and
- *not affected* by $|i|$ in full disclosure equilibrium.

Road Map

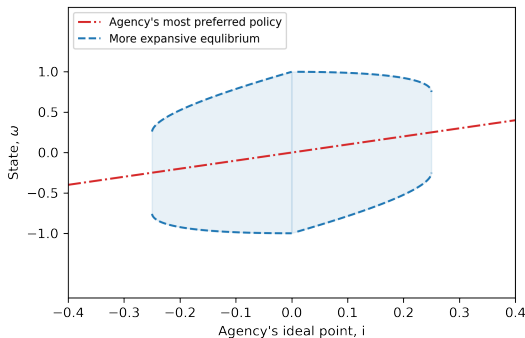
- ① Introduction
- ② Uniform Prior, Quadratic Preferences, State-independent Sender Preferences
 - Game Structure
 - Equilibrium Characterization
 - Effects of Preference Divergence
 - Belief Stability Equilibrium Refinement
- ③ General Model
 - Equilibria
 - Comparative Statics
- ④ Robustness
- ⑤ Summary

Belief-Stability: Motivation

We have multiple equilibria with contrary comparative statics:

- More expansive $\rightarrow$ communication deteriorates in ex-ante preference misalignment
- Less expansive $\rightarrow$ communication improves in ex-ante preference misalignment

All survive standard refinements $\rightarrow$ Which one should we expect?

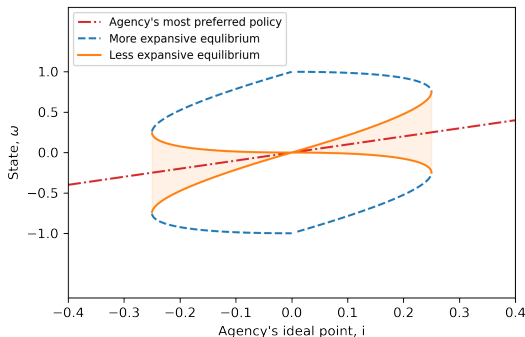


Belief-Stability: Motivation

We have multiple equilibria with contrary comparative statics:

- More expansive $\rightarrow$ communication deteriorates in ex-ante preference misalignment
- Less expansive $\rightarrow$ communication improves in ex-ante preference misalignment

All survive standard refinements $\rightarrow$ Which one should we expect?



Belief-Stable Equilibria

Definition: Belief Stable

Consider an equilibrium (σ, μ) .

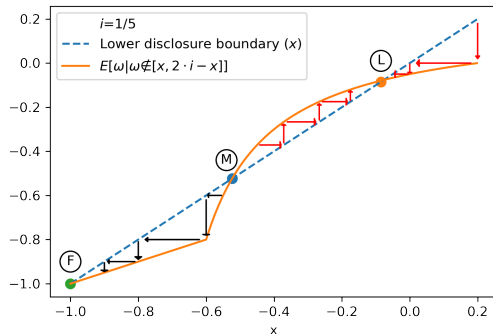
Let μ_j^ε be j 's perturbed system of beliefs.

Let σ^ε be seq. rational given $(\mu_j^\varepsilon, \mu_{-j})$.

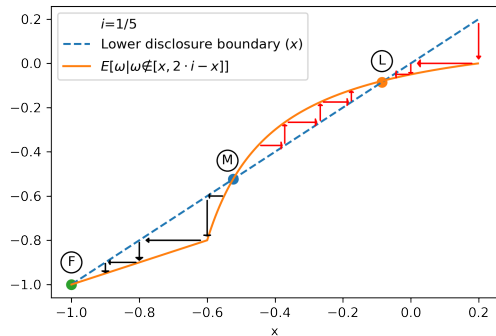
Let $\hat{\mu}_j^\varepsilon$ be consistent with σ^ε .

If there exists an $\varepsilon > 0$ such that, for every μ_j^ε that satisfies $|\mu_j^\varepsilon(y) - \mu_j(y)| < \varepsilon$, $|\hat{\mu}_j^\varepsilon(y) - \mu_j(y)| \leq |\mu_j^\varepsilon(y) - \mu_j(y)|$ is satisfied for every decision node y of j

$\Rightarrow$ Equilibrium (σ, μ) is **belief-stable** (for j)



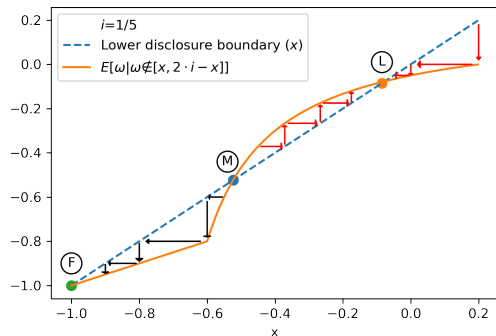
Belief-Stable Equilibria



Prop.3

- 1 More expansive equilibrium is not belief-stable

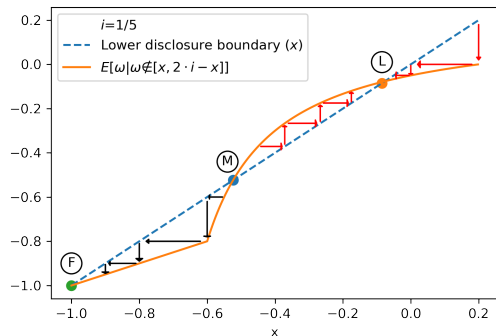
Belief-Stable Equilibria



Prop.3

- ① More expansive equilibrium is not belief-stable;
- ② Less expansive equilibrium is belief-stable;

Belief-Stable Equilibria

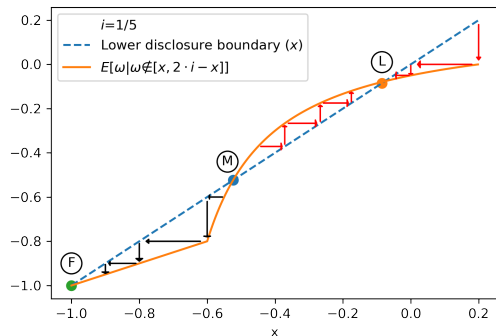


Prop.3

- ① More expansive equilibrium is not belief-stable;
- ② Less expansive equilibrium is belief-stable;

⇒ **Corollary 1.** Equilibrium is belief-stable $\Leftrightarrow$ communication **improves** in pref. divergence.
Equilibrium is not belief-stable $\Leftrightarrow$ communication **worsens** in pref. divergence.

Belief-Stable Equilibria



Prop.3

- ① More expansive equilibrium is not belief-stable;
- ② Less expansive equilibrium is belief-stable;
- ③ Full disclosure is belief-stable when $i \neq 0$.

⇒ **Corollary 1.** Equilibrium is belief-stable $\Leftrightarrow$ communication **improves** in pref. divergence.
Equilibrium is not belief-stable $\Leftrightarrow$ communication **worsens** in pref. divergence.

Road Map

- ① Introduction
- ② Uniform Prior, Quadratic Preferences, State-independent Sender Preferences
 - Game Structure
 - Equilibrium Characterization
 - Effects of Preference Divergence
 - Belief Stability Equilibrium Refinement
- ③ General Model
 - Equilibria
 - Comparative Statics
- ④ Robustness
- ⑤ Summary

General Model: Actors and Timing

Two players: the Agency (it) and the Policymaker (she).

①	Nature determines state of the world $\omega \in \Omega$: Ω is compact and $\text{conv}(\Omega) = [\underline{\omega}, \bar{\omega}]$	$\omega \sim F(\cdot)$ such that $\int_{\underline{\omega}}^{\bar{\omega}} x \cdot f(x) dx = 0$
②	Agency observes ω	ω
③	Agency chooses message (m) to send to Policymaker	$m \in \{\omega, \emptyset\}$
④	Policymaker observes m and chooses policy (p) to implement	$p \in \mathbb{R}$

$$p^P(\omega) := \arg \max_p u_P(p; \omega) = \omega$$

$$p^A(\omega, \alpha, i) := \arg \max_p u_A(p; \omega, \alpha, i) = \alpha \cdot p^P(\omega) + (1 - \alpha) \cdot i$$

Full-Disclosure Equilibria

Proposition

- ① If $u_A(p^P(\bar{\omega}); \bar{\omega}, \alpha, i) > u_A(p^P(\underline{\omega}); \bar{\omega}, \alpha, i)$, then for all conditions on primitives such that there exists a full-disclosure equilibrium s.t. $x^* = p^P(\underline{\omega})$, that equilibrium is belief-stable;
 - ② If $u_A(p^P(\bar{\omega}); \bar{\omega}, \alpha, i) = u_A(p^P(\underline{\omega}); \bar{\omega}, \alpha, i)$, then for all conditions on primitives, such that there exists a full-disclosure equilibrium s.t. $x^* = p^P(\underline{\omega})$, that equilibrium is belief-unstable;
 - ③ If $u_A(p^P(\bar{\omega}); \bar{\omega}, \alpha, i) < u_A(p^P(\underline{\omega}); \bar{\omega}, \alpha, i)$, there does not exist a full-disclosure equilibrium such that $x^* = p^P(\underline{\omega})$.
- ② (Seidman and Winter 1997) There exists a full-disclosure equilibrium if the Agency's utility, $u_A(\cdot)$, satisfies single-crossing.

Non-Disclosure Equilibria

Proposition

- ① There exists a unique threshold $\alpha^* \in (0, 1)$ such that a non-disclosure equilibrium exists if and only if the Agency's bias $i = p_0^P$ and preference state-dependence $\alpha \leq \alpha^*$.
- ② A non-disclosure equilibrium is belief-stable.

Partial-Disclosure Equilibria

Proposition

There exists a threshold $\alpha^{**} \in [\alpha^*, 1)$ and, $\forall \alpha \leq \alpha^{**}$, an interval $I^*(\alpha) \subset \Omega$ containing p_0^P such that $G \in \mathcal{G}$ has a partial-disclosure equilibrium if and only if $\alpha \leq \alpha^{**}$ and $i \in I^*(\alpha)$.

Equilibrium Pattern

Proposition

- ① For any distinct x_j^* and x_k^* ,

$$M(x_k^*, \alpha, i) \subseteq M(x_j^*, \alpha, i) \Leftrightarrow u_A(x_k^*; \omega, \alpha, i) \geq u_A(x_j^*; \omega, \alpha, i).$$

- ② Index set $X^*(\alpha, i)$ s.t. if $j < k$, $x_j^* < x_k^*$. Then the belief-stability of equilibria alternates along this ordering, i.e., if the equilibrium corresponding to x_j^* is belief-stable, then the equilibrium corresponding to x_{j+1}^* (if it exists) is not belief-stable and the equilibrium corresponding to x_{j+2}^* (if it exists) is.

Corollary

For any given (α, i) , knowing the belief-stability of one equilibrium is sufficient to determine the belief-stability of all equilibria.

General Model: Comparative Statics

Proposition

The equilibrium policy absent disclosure, x^* , is

- ① decreasing in the Agency's bias, i ; and
- ② increasing in the Agency's preference state-dependence, α , when $x^* < i$, and decreasing in α otherwise,

if and only if the equilibrium is belief-stable.

Proposition

The equilibrium disclosure interval, $M(x^*, \alpha, i)$, is

- ① expanding in the Agency's bias, i , when $x^* \leq i$, and contracting in i otherwise; and
- ② expanding in the Agency's preference state-dependence, α

if and only if the equilibrium is belief-stable.

Comparative Statics Cont'd

When is $x^* < i$?

Remark

When

- ① distance-based utilities,
- ② α not too large, and
- ③ the tails of the prior density $f(\omega)$ are not too asymmetric

Road Map

- ① Introduction
- ② Uniform Prior, Quadratic Preferences, State-independent Sender Preferences
 - Game Structure
 - Equilibrium Characterization
 - Effects of Preference Divergence
 - Belief Stability Equilibrium Refinement
- ③ General Model
 - Equilibria
 - Comparative Statics
- ④ Robustness
- ⑤ Summary

Robustness

- Messages are Partially Verifiable
 - retain equilibrium partial disclosure
 - disclosure increases in i
- Vague Messages
 - retain equilibrium partial disclosure
 - disclosure increases in i

Road Map

- ① Introduction
- ② Uniform Prior, Quadratic Preferences, State-independent Sender Preferences
 - Game Structure
 - Equilibrium Characterization
 - Effects of Preference Divergence
 - Belief Stability Equilibrium Refinement
- ③ General Model
 - Equilibria
 - Comparative Statics
- ④ Robustness
- ⑤ Summary

Summary

A model of **verifiable communication** between a Policymaker and a Bureaucratic Agency:

- ① When Agency and Policymaker's ex-ante preferences are sufficiently aligned, unraveling may stop before being complete;
- ② Greater ex-ante preference divergence can encourage Agency to disclose more information;
- ③ Equilibria where communication improves with preference divergence are belief-stable.

Thank you!