



The optimality of (stochastic) veto delegation [☆]

Xiaoxiao Hu ^a, Haoran Lei ^{b,*}

^a Business School, Ningbo University, China

^b School of Economics and Trade, Hunan University, China

ARTICLE INFO

JEL classification:

D72

D82

Keywords:

Optimal delegation

Veto delegation

Stochastic mechanism

ABSTRACT

We analyze the optimal delegation problem between a principal and an agent, assuming that the latter has state-independent preferences. We demonstrate that if the principal is more risk-averse than the agent toward non-status quo options, an optimal mechanism is a *veto mechanism*. In a veto mechanism, the principal uses veto (i.e., maintaining the status quo) to balance the agent's incentives and does not randomize among non-status quo options. We characterize the optimal veto mechanism in a one-dimensional setting. In the solution, the principal uses veto only when the state surpasses a critical threshold.

Contents

1. Introduction	2
1.1. Motivating example	2
1.2. Literature	3
2. Model	3
2.1. Veto mechanisms	4
2.2. Remarks on the status quo option	5
3. Characterization	5
3.1. Analysis	5
3.2. When are veto mechanisms valuable?	7
3.3. Comparative statics	8
4. Discussions	10
4.1. Scenario 1: $v_0 \geq 1$	10
4.2. Scenario 2: $v_0 \in (0, 1)$	10
5. Concluding remarks	11
Declaration of competing interest	11
Appendix A.	11
A.1. Proof of Proposition 2	12
A.2. Proof of Corollary 1	14
A.3. Proof of Proposition 3	15
A.4. Proof of Proposition 4	15
A.5. Veto probability $\bar{p}^*(\theta)$ may not change uniformly as v_0 varies	15

[☆] We thank Navin Kartik and participants at various seminars and conferences for helpful comments. All errors are ours.

* Corresponding author.

E-mail addresses: xhuah@connect.ust.hk (X. Hu), hleiaa@connect.ust.hk (H. Lei).

A.6.	Insufficiency of the likelihood ratio conditions	16
A.7.	Proof of Proposition 5	16
A.8.	Derivations omitted in Section 4.2	17
A.8.1.	U-shaped $\tilde{a}^*(\theta)$ and $\tilde{p}^*(\theta)$	17
A.8.2.	Deterministic optimal mechanism	19
Data availability		19
References		19

1. Introduction

In many settings of economic and political interest, a principal consults an informed but biased agent while contingent transfers between them are infeasible. The principal then specifies a permissible set of options, termed as the *delegation set*, and allows the agent to choose any option from this set. Holmstrom (1980) analyzes the delegation problem when the agent's utility depends on both the chosen option and the realized state. He characterizes the optimal delegation sets, focusing on the case when the delegated options form an interval (i.e., *interval delegation*). There is a large literature following Holmstrom (1980),¹ most of which assume state-dependent agent preferences and emphasize the optimality of interval delegation.

Nevertheless, in many real-life interactions, agents have *state-independent* preferences and only care about the final decisions. For example, prosecutors may seek their own favorable rulings regardless of the severity of defendants' crimes, and financial advisors may recommend higher-commission products irrespective of their appropriateness. When the agent has state-independent preferences, any interval delegation set fails to elicit the agent's private information and the principal cannot leverage the agent's expertise through interval delegation. On the other hand, interval delegation also seems incompatible with the *veto delegation* ubiquitous in many organizations, where the principal can only either accept the agent's proposal or reject it in favor of an outside option.

Veto delegation has been widely used in various organizations, especially in the legislative processes and corporate governance. In the legislation, veto delegation is known as the *closed rule*, where constituents can either approve or veto a committee's proposed bill. The closed rule is deemed to constitute a "critical component of managerial power in the U.S. House of Representatives" (Doran, 2010). In corporate governance, boards of directors often can only disapprove proposals but cannot unilaterally enact new ones. More instances of the veto delegation in other organizations are documented in Marino (2007), Mylovanov (2008) and Lubensky and Schmidbauer (2018).

Motivated by these observations, we analyze the optimal delegation problem between a principal and an agent where the agent's preference is state-independent. In our model, the principal decides whether to maintain the status quo or choose another option. To elicit information from the biased agent, the principal commits to a (possibly stochastic) direct mechanism. We demonstrate that a principal's optimal mechanism must be a *veto mechanism*, provided that the principal is more risk-averse than the agent regarding non-status quo options. The veto mechanism is a specific direct mechanism in which the principal never randomizes among non-status quo options and only uses veto to balance the agent's incentives for truth-telling. Our result rationalizes veto delegation as an arrangement that achieves the optimal outcome.

1.1. Motivating example

To illustrate our main result, consider a political advisor (agent) advising a policy maker (principal).² There are two equally likely states, θ_1 and θ_2 , and the agent privately observes the realized one. The policy maker decides between a new policy (a_1 or a_2) or maintaining the status quo policy a_0 . The mild policy a_1 is suitable when the state is θ_1 , and the aggressive policy a_2 is suitable when the state is θ_2 . The principal obtains a payoff of 1 if the enacted new policy is suitable, or -3 otherwise; the status quo a_0 yields a constant payoff of 0 for the principal. The advisor's preference ranks a_2 highest and a_0 lowest, regardless of the state. Suppose the advisor's payoff is i when a_i is chosen for $i \in \{0, 1, 2\}$. Table 1 shows players' payoffs for each policy–state combination.

The policy maker's ex ante optimal choice is a_0 , yielding a payoff of 0 for both players. Can the policy maker elicit information from the extremely biased advisor through delegation? The answer is yes. The policy maker can commit to the following mechanism where she may accept the advisor's proposal or *use veto* (i.e., choosing a_0): if the advisor proposes a_1 , the policy maker will approve; if a_2 is proposed, the policy maker will approve or veto with equal probabilities. Given the principal's commitment, the advisor is indifferent between proposing either new policy. Suppose the advisor proposes a_i at each state θ_i . Then, the policy maker's expected utility is $3/4 > 0$. The described veto mechanism can be viewed as the policy maker delegating two *lotteries* to the advisor: $l_1 = (1 \circ a_1)$ and $l_2 = (0.5 \circ a_0, 0.5 \circ a_2)$.³ Therefore, the policy maker benefits from delegating his decisions to the agent when stochastic mechanisms are allowed.

One can verify that the described mechanism, while being conceptually simple and relatively easy to implement, yields the highest payoff for the principal among all incentive-compatible mechanisms. The key feature of the veto mechanism is that the principal only

¹ We briefly review the literature on optimal delegation in Section 1.2.

² This example is adapted from the think-tank game in Lipnowski and Ravid (2020), pp. 1632–1634.

³ Throughout this paper we denote a lottery with finite possible outcomes by $l = (p_1 \circ a_1, \dots, p_n \circ a_n)$ with $\sum_{i=1}^n p_i = 1$, meaning that outcome a_i occurs with probability p_i for each $i \in \{1, \dots, n\}$.

Table 1
Payoffs of the policy maker and the political advisor.

	a_0	a_1	a_2
θ_1	(0, 0)	(1, 1)	(−3, 2)
θ_2	(0, 0)	(−3, 1)	(1, 2)

uses veto to balance the agent's incentives; that is, she never randomizes between non-status quo options. Later, we demonstrate the optimality of the veto mechanism in a general environment. In this sense, our paper contributes to understanding why veto delegation is prevalent in various organizations.

1.2. Literature

Our paper is mostly related to the literature on optimal delegation. Starting from Holmstrom (1980), seminal works in this literature include Melumad and Shibano (1991), Alonso and Matouschek (2008), Kováč and Mylovanov (2009), Amador and Bagwell (2013), Kolotilin and Zapechelnyuk (2019), etc. Variants of the delegation model have been used in political economy (Bendor et al., 2001; Krishna and Morgan, 2001), monopoly regulation and organization (Baron and Myerson, 1982; Aghion and Tirole, 1997), tariff policies (Amador and Bagwell, 2013), etc. Despite the vast body of existing literature, this paper is, to the best of our knowledge, the first to explore the one-shot⁴ optimal delegation problem where the agent has state-independent preferences. State-independent agent preferences imply that deterministic mechanisms are generally suboptimal. Specifically, if the agent's utility function is injective, no incentive-compatible deterministic mechanisms can elicit the agent's private information. The potential non-optimality of deterministic mechanisms in delegation has been demonstrated in previous studies, such as Kováč and Mylovanov (2009, Section 4) and Kartik et al. (2021, Appendix E). In the same paper, Kartik, Kleiner, and Van Weelden also provide a practical real-life example of using stochastic delegation in nominating judges.

As we interpret our result as demonstrating the optimality of veto delegation, we discuss the relationship between this paper and the growing literature on veto delegation (Dessein, 2002; Marino, 2007; Mylovanov, 2008; Lubensky and Schmidbauer, 2018). The above-mentioned works differ from the optimal delegation literature as they focus on specific game forms of veto delegation—the agent proposes some option and the principal either accepts it or rejects it for an outside option. Fixing an exogenous outside option, Dessein (2002) compares full delegation and veto delegation and shows that full delegation dominates veto delegation as long as the conflict of interest is not extreme. Lubensky and Schmidbauer (2018) strengthen the result of Dessein (2002), demonstrating that full delegation is superior by explicitly characterizing the most informative veto equilibrium. Marino (2007) challenges Dessein's conclusion and shows the superiority of veto delegation in another setting with different assumptions of players' preferences and prior distribution. Mylovanov (2008) allows the outside option to be chosen endogenously and shows the equivalence of optimal delegation and optimal veto delegation. It is worth emphasizing that unlike those works mentioned above, we adopt a mechanism design approach and the veto delegation arrangement emerges as the principal's optimal mechanism. Another important distinction is that in the above-mentioned works, the principal decides whether to veto or not after the agent's proposal; however, in our paper the principal commits ex ante to his probability of veto conditional on the agent's proposal to ensure information provision from the extremely biased agent.

Finally, our paper relates to the broader information transmission literature. The assumption of sender (i.e., agent) state-independent preferences is common in the literature of communication with hard evidence (Glazer and Rubinstein, 2004, 2006; Hart et al., 2017) and information design (e.g., the judge-prosecutor game in Kamenica and Gentzkow (2011)). There is also a growing literature on cheap talk with sender state-independent preferences (Chakraborty and Harbaugh, 2010; Lipnowski and Ravid, 2020; Diehl and Kuzmics, 2021). While this assumption is arguably of great empirical relevance, it has rarely been studied in the existing delegation literature. This paper therefore fills the gap.

2. Model

There are two players, a principal (he) and a better-informed agent (she). The state θ follows some full-support distribution $\mu \in \Delta(\Theta)$, where the state space Θ is a subset of $\mathbb{R}^n$ for some positive integer n . The agent privately observes the realized state. Denote by $a \in A$ a generic option or allocation, where the set of available options A is a subset of $\mathbb{R}^\ell$ for some positive integer ℓ . Principal's and agent's payoff functions are $u : A \times \Theta \rightarrow \mathbb{R}$ and $v : A \rightarrow \mathbb{R}$, respectively. While the principal's payoffs depend on the realized state, the agent's do not.

Status-quo option Notably, there exists a unique *status quo option* $a_0 \in A$, and any option distinct from a_0 is referred as a *non-status quo option*. We interpret choosing a_0 as maintaining the status quo option or allocation. Real-life instances of a_0 include customers buying nothing from the intermediaries, investors rejecting the financial products recommended by the financial advisors, and policy makers vetoing the political advisors' proposals for new policies. The principal is assumed to be *more risk-averse* than the agent toward choosing non-status quo options at *all* states. Formally, for each $\theta \in \Theta$, there exists a strictly concave transformation $h_\theta : \mathbb{R} \rightarrow \mathbb{R}$ such that $u(a, \theta) = h_\theta(v(a))$ for all $a \neq a_0$. Additionally, the set $v(A \setminus \{a_0\}) = \{v(a) : a \in A \text{ and } a \neq a_0\}$ is assumed to be connected.

⁴ Frankel (2016) and Chen (2022) study the dynamic delegation problem with agent state-independent preferences.

Following the literature on optimal delegation (e.g., Melumad and Shibano, 1991; Martimort and Semenov, 2006; Alonso and Matouschek, 2008), we approach the delegation problem from a mechanism design perspective and focus on the principal's optimal mechanism. A *direct mechanism* is a Borel measurable function $m : \Theta \rightarrow \Delta(A)$ such that expected payoffs are integrable. To simplify notations, we extend the domains of players' payoff functions to incorporate stochastic options:

$$u(\alpha, \theta) := \int_A u(a, \theta) d\alpha \quad \text{and} \quad v(\alpha) := \int_A v(a) d\alpha$$

where $\alpha \in \Delta(A)$ denotes some stochastic option. A mechanism m is *incentive compatible* if

$$v(m(\theta)) \geq v(m(\theta')) \text{ for all } \theta, \theta' \in \Theta. \quad (\text{IC})$$

Constraints (IC) can be reduced to a set of indifference constraints:

$$v(m(\theta)) = \hat{v} \text{ for all } \theta \in \Theta \text{ and some } \hat{v} \in \mathbb{R}. \quad (\text{I})$$

As the agent's preference is state-independent, the reduced incentive-compatibility constraints (I) are also state-independent. That is, in an incentive-compatible mechanism, the agent obtains the same utility from any report regardless of the state. Principal's maximization problem is given by:

$$\max_{m, \hat{v} \in \mathbb{R}} \int_{\Theta} u(m(\theta), \theta) d\mu \text{ subject to constraints (I)}. \quad (\mathcal{M})$$

A direct mechanism m^* is called an *optimal mechanism* if $(m^*, \hat{v}^*)$ solves the problem (M) where $\hat{v}^* = v(m^*(\theta))$ for all θ . The solution concept is *perfect Bayesian equilibrium* (henceforth, *equilibrium*). One interpretation of our mechanism design approach, as in most applied mechanism design papers, is to find an upper bound on the principal's welfare. Having said that, we also find the implementation via veto delegation fitting naturally into various contexts.

2.1. Veto mechanisms

A direct mechanism m is a *veto mechanism* if there exist two mappings, $\bar{p} : \Theta \rightarrow [0, 1]$ and $\bar{a} : \Theta \rightarrow A \setminus \{a_0\}$, such that for almost every state θ with respect to the prior distribution μ , the induced stochastic option $m(\theta)$ assigns probability $\bar{p}(\theta)$ to the status quo option a_0 and the complementary probability $1 - \bar{p}(\theta)$ to the non-status quo option $\bar{a}(\theta)$. This definition corresponds to the real-life veto delegation scenario where the agent proposes some option $\bar{a}(\theta)$ and then the principal vetoes that proposal with probability $\bar{p}(\theta)$. We refer to $\bar{p}(\theta)$ as the *veto probability*. The veto mechanism encompasses all deterministic mechanisms as well as a class of simple stochastic mechanisms. In these stochastic mechanisms, whenever the principal makes a random choice, the randomization occurs only between the status quo and one other option.

Proposition 1 allows us to focus on veto mechanisms when searching for an optimal mechanism.

Proposition 1. *A direct mechanism is optimal only if it is a veto mechanism.*

Proof. Principal's maximization problem (M) can be solved sequentially. First, fix some agent utility $\hat{v}$ and solve the following optimization problem

$$\pi(\hat{v}) := \max_m \int_{\Theta} u(m(\theta), \theta) d\mu \text{ subject to } v(m(\theta)) = \hat{v} \text{ for all } \theta. \quad (\mathcal{M}_{\hat{v}})$$

Then, solve for $v^* \in \arg \max_{\hat{v} \in \mathbb{R}} \pi(\hat{v})$ that maximizes principal's ex ante payoff. For our purpose here, it suffices to show that for all $\hat{v} \in v(A)$, any direct mechanism $\hat{m}$ that solves the problem $(\mathcal{M}_{\hat{v}})$ must be a veto mechanism.

Fixing some $\hat{v} \in v(A)$, mechanism $\hat{m}$ is a solution to $(\mathcal{M}_{\hat{v}})$ only if, for μ -almost all θ , $\hat{m}(\theta)$ solves the following maximization problem⁵:

$$\max_{\alpha \in \Delta(A)} u(\alpha, \theta) \text{ subject to } v(\alpha) = \hat{v}. \quad (\mathcal{M}_{\hat{v}, \theta})$$

Fixing some state θ , the objective function in $(\mathcal{M}_{\hat{v}, \theta})$ can be written as

$$u(\alpha, \theta) = \int_{A \setminus \{a_0\}} u(a, \theta) d\alpha + p_0 u(a_0, \theta)$$

where p_0 is the probability assigned to a_0 under α . Further,

⁵ Precisely, $\hat{m}$ is a solution to maximization problem $(\mathcal{M}_{\hat{v}})$ if and only if $v(\hat{m}(\theta)) = \hat{v}$ for all $\theta \in \Theta$ and there exists some subset $\hat{\Theta}$ of Θ with $\mu(\hat{\Theta}) = 1$ such that $\hat{m}$ solves $(\mathcal{M}_{\hat{v}, \theta})$ for all $\theta \in \hat{\Theta}$.

$$\begin{aligned}
\int_{A \setminus \{a_0\}} u(a, \theta) d\alpha &= (1 - p_0) \int_{A \setminus \{a_0\}} h_\theta(v(a)) d\left(\frac{\alpha}{1 - p_0}\right) \\
&\leq (1 - p_0) h_\theta\left(\int_{A \setminus \{a_0\}} v(a) d\left(\frac{\alpha}{1 - p_0}\right)\right) \quad (\text{because } h_\theta \text{ is concave}) \\
&= (1 - p_0) h_\theta(v(a')) \text{ for some } a' \in A \setminus \{a_0\} \quad (\text{because } v(A \setminus \{a_0\}) \text{ is connected}) \\
&= (1 - p_0) u(a', \theta).
\end{aligned}$$

Since function h_θ is strictly concave, the inequality takes equality if and only if the support of α contains at most one option distinct from a_0 . Therefore, any solution to the problem $(\mathcal{M}_\theta)$ must be a veto mechanism. $\square$

The optimality of veto mechanisms relies on the assumption that the principal is more risk-averse than the agent towards non-status quo options. If the principal is less risk-averse than the agent, the optimality of veto mechanisms may fail. Below, we provide an example to illustrate this scenario. Suppose $\Theta = [0, 1] \subseteq \mathbb{R}$ and $A = (-\infty, 1] \cup \{a_0\} \subseteq \mathbb{R}$, where the status quo option is denoted by some real number $a_0 > 1$. The principal has absolute loss payoff: $u(a, \theta) = -|a - \theta|$ for $a \in \mathbb{R}$ and $u(a_0, \theta) = u_0 < 0$ for all θ . The agent has quadratic loss payoff $v(a) = -(1 - a)^2$ for $a \in (-\infty, 1]$ and $v(a_0) = v_0 \in \mathbb{R}$. Then, the principal can achieve an outcome arbitrarily close to the first-best outcome without using veto at all.⁶ For arbitrarily small $\varepsilon > 0$, given the reported state θ the principal chooses a lottery of actions with expectation being $\theta - \varepsilon$, its support lying in $(-\infty, \theta]$ and the appropriate variance such that constraints (I) hold.

2.2. Remarks on the status quo option

The proof of Proposition 1 can be extended to show that if the principal is more risk-averse than the agent towards *all* options at all states (and $v(A)$ is connected), an optimal mechanism must be deterministic. In this case, the principal generally cannot elicit any information from the agent due to the state-independent agent preference. Therefore, for the delegation problem to be non-trivial, despite the fact that the principal is more risk-averse than the agent towards all non-status quo options at all states, adding the status quo option will make the principal no longer more risk-averse than the agent. In other words, some player must have different risk attitudes between the status quo option and non-status quo options.

In behavioral economics, it is well-documented that individuals perceive qualitative differences between choosing the status quo option and a non-status quo option (e.g., Kahneman and Tversky, 1979). Additionally, individuals often feel more responsible when enacting a non-status quo option (Bartling and Fischbacher, 2012). In our motivating example, the policy maker is insensitive to the outcome of maintaining the status quo in that choosing a_0 yields the same payoff at both states. One possible reason is that the policy maker feels more responsible for the outcome of a non-status quo policy than that of maintaining the status quo.

3. Characterization

We characterize the optimal mechanism in a one-dimensional setting. Assume $\Theta = [0, 1] \subseteq \mathbb{R}$ and $A = \{a_0\} \cup [-M, M] \subset \mathbb{R}$, where $M (> 1)$ is a sufficiently large positive number and the status quo option is denoted by some real number $a_0 < -M$. The state θ follows some full-support distribution $\mu \in \Delta(\Theta)$. Agent's payoff function is $v(a) = a$ for $a \in [-M, M]$ and $v(a_0) = v_0 \in \mathbb{R}$. Principal's payoff function is $u(a, \theta) = -(a - \theta)^2$ for $a \in [-M, M]$ and $u(a_0, \theta) = u_0$ for all $\theta \in \Theta$. While the principal's payoffs from a non-status quo option depend on the state, those from the status quo option do not.⁷

We impose the following restrictions. Assume $u_0 \leq 0$. Otherwise, the status quo option a_0 always yields the unique highest payoff for the principal among all options, and the principal trivially chooses a_0 regardless of the state. We also assume $v_0 \leq 0$ when characterizing the optimal mechanism and discussing comparative statics. Later, we show that allowing v_0 being positive can drastically affect the optimal mechanism (Section 4).

3.1. Analysis

We first briefly discuss the extreme case when $u_0 = 0$. In this case, the principal can never do better than opting for the status quo option, and there exists a *pooling equilibrium* in which he maintains the status quo regardless of the state. Nevertheless, there also exists a *fully separating equilibrium*, in which the agent is strictly better off. In the separating equilibrium, the principal commits to the

⁶ This almost first-best outcome is essentially the same with the example of non-optimality of deterministic allocations in Kováč and Mylovánov (2009, Section 4). A related example can be found in Example E.1 in Kartik et al. (2021).

⁷ The assumption that principal's payoffs from a_0 are state-independent aligns with many real-life scenarios. For instance, in a customer-salesperson interaction, the customer "delegates" his purchasing decision to the salesperson, and the state captures the degrees to which the available products match the customer's needs. As the value of the alternative use of money for the customer does not depend on how well the products suit her, the payoffs from rejecting the salesperson's recommendation are state-independent. In the literature of contract theory, it is also usually assumed that the value of outside option from not signing the contract is state-independent for the principal. For example, in an investor-entrepreneur relationship, the participation constraint requires that the expected return to the investor (principal) must cover the opportunity cost of funds, and the opportunity cost is usually measured by a constant interest rate (e.g., Gale and Hellwig, 1985).

veto mechanism $(\tilde{a}^*, \tilde{p}^*)$ where $\tilde{a}^*(\theta) = \theta$ and $\tilde{p}^*(\theta) = \theta/(\theta - v_0)$. That is, given the agent's reported state θ , the principal uses veto with probability $\frac{\theta}{\theta - v_0}$ and chooses option $\tilde{a}(\theta) = \theta$ with the complementary probability. This mechanism is incentive-compatible, and the principal's ex ante payoff is the same as that of the pooling equilibrium. The separating equilibrium is the limit of equilibria with $u_0 < 0$, as discussed later.

From now on, assume $u_0 < 0$. By Proposition 1, we focus on veto mechanisms and the principal's maximization problem $(\mathcal{M})$ reduces to

$$\begin{aligned} & \max_{\{\tilde{p}(\theta), \tilde{a}(\theta)\}_{\theta \in \Theta, \hat{v}}} \int_{\Theta} (\tilde{p}(\theta)u_0 - (1 - \tilde{p}(\theta))(\tilde{a}(\theta) - \theta)^2) d\mu \\ & \text{subject to (i) } \forall \theta \in \Theta, \tilde{p}(\theta)v_0 + (1 - \tilde{p}(\theta))\tilde{a}(\theta) = \hat{v}; \\ & \quad \text{(ii) } \forall \theta \in \Theta, \tilde{p}(\theta) \in [0, 1] \text{ and } \tilde{a}(\theta) \in [-M, M]. \end{aligned} \quad (\mathcal{M}')$$

Perturbing $\tilde{p}(\theta')$ and $\tilde{a}(\theta')$ at some state θ' will not impact the validity of constraints (I) for other states $\theta \neq \theta'$, as long as the agent's payoff remains fixed at $\hat{v}$ when reporting θ' . Due to this observation, we first fix some agent utility $\hat{v}$ and solve the principal's maximization problem *state-wisely*. Then, we solve for the optimal $\hat{v}$.

The original problem $(\mathcal{M}')$ is reduced to the following two-step maximization problem:

1. For each $\theta \in \Theta$ and each agent's utility $\hat{v} \in [v_0, 1]$, solve the following optimization problem:

$$\begin{aligned} \pi(\hat{v}, \theta) &\equiv \max_{\tilde{p}(\theta), \tilde{a}(\theta)} \tilde{p}(\theta)u_0 - (1 - \tilde{p}(\theta))(\tilde{a}(\theta) - \theta)^2 \\ & \text{subject to (i) } \tilde{p}(\theta)v_0 + (1 - \tilde{p}(\theta))\tilde{a}(\theta) = \hat{v}; \\ & \quad \text{(ii) } \tilde{p}(\theta) \in [0, 1] \text{ and } \tilde{a}(\theta) \in [-M, M]. \end{aligned} \quad (\mathcal{M}_1)$$

2. Solve for the optimal $\hat{v}$:

$$\max_{\hat{v} \in [v_0, 1]} \int_{\Theta} \pi(\hat{v}, \theta) d\mu. \quad (\mathcal{M}_2)$$

We solve the sequential problems $(\mathcal{M}_1)$ and $(\mathcal{M}_2)$ using the standard Lagrangian method. The optimal veto mechanism is summarized in Proposition 2.

Proposition 2. Assume $v_0 \leq 0$.

- (a) If $(1 - v_0)^2 + u_0 > (\mathbb{E}_{\mu}(\theta) - v_0)^2$, then define $\eta(\theta) \equiv \sqrt{(\theta - v_0)^2 + u_0} + v_0$ for $\theta \geq \sqrt{-u_0} + v_0$ and the optimal veto mechanism is

$$\tilde{a}^*(\theta) = \begin{cases} \bar{a}, & \text{if } \theta < \bar{\theta}; \\ \eta(\theta), & \text{if } \theta \geq \bar{\theta}, \end{cases} \quad \tilde{p}^*(\theta) = \begin{cases} 0, & \text{if } \theta < \bar{\theta}; \\ \frac{\eta(\theta) - \bar{a}}{\eta(\theta) - v_0}, & \text{if } \theta \geq \bar{\theta}, \end{cases} \quad (1)$$

where $\bar{\theta} \equiv \eta^{-1}(\bar{a})$ and $\bar{a}$ is determined by $\int_{\eta^{-1}(\bar{a})}^1 (\eta(\theta) - \bar{a}) d\mu(\theta) + \bar{a} - \mathbb{E}_{\mu}(\theta) = 0$. Moreover, $\bar{a} \in (0, \mathbb{E}_{\mu}(\theta))$.

- (b) If $(1 - v_0)^2 + u_0 \leq (\mathbb{E}_{\mu}(\theta) - v_0)^2$, then the optimal veto mechanism is characterized by $\tilde{a}^*(\theta) = \mathbb{E}_{\mu}(\theta)$ and $\tilde{p}^*(\theta) = 0$ for all $\theta \in \Theta$; that is, the principal chooses $\mathbb{E}_{\mu}(\theta)$ regardless of the state.

Proof. See Appendix A.1. $\square$

Since the principal can always choose his ex ante favorite option by trivially delegating a singleton set, he strictly benefits from eliciting the agent's private information when a non-trivial mechanism is optimal, as in Case (a) of Proposition 2. Fig. 1 illustrates the optimal veto mechanism in this case: the principal's choices are pooled at the deterministic action $\bar{a}$ when the state is below the threshold $\bar{\theta}$; when the state is above the threshold $\bar{\theta}$, the principal uses veto with probability $\frac{\eta(\theta) - \bar{a}}{\eta(\theta) - v_0}$ and chooses $\eta(\theta)$ with the complementary probability. In Case (b), the principal trivially delegates $\{\mathbb{E}_{\mu}(\theta)\}$.

One interpretation of the non-trivial optimal veto mechanism is through delegating lotteries. The principal delegates a set of lotteries $\{l_{\theta}\}_{\theta \in [\bar{\theta}, 1]}$ to the agent, indexed by θ , where $l_{\theta} = (\tilde{p}^*(\theta) \circ a_0, (1 - \tilde{p}^*(\theta)) \circ \eta(\theta))$ for $\theta \in [\bar{\theta}, 1]$. The agent chooses the lottery l_{θ} at $\theta \in [\bar{\theta}, 1]$, and chooses $l_{\bar{\theta}}$ otherwise. Another interpretation is through the veto delegation arrangement, where the agent proposes an option and the principal either vetoes or approves. Specifically, the principal pre-commits to the following stochastic choice contingent on the agent's proposal: the principal uses veto with probability $\tilde{p}^*(\eta^{-1}(\hat{a}))$ if the proposed option $\hat{a} \in [\bar{a}, \eta(1)]$, always approves if $\hat{a} \in [0, \bar{a}]$ and otherwise always uses veto.

Fig. 1 suggests that both $\tilde{a}^*(\theta)$ and $\tilde{p}^*(\theta)$ are increasing over $[\bar{\theta}, 1]$ and, perhaps surprisingly, that action $\tilde{a}^*(\theta)$ is strictly lower than θ for all $\theta > \bar{\theta}$. We summarize these properties in Corollary 1, which can be verified directly from the corresponding expressions.

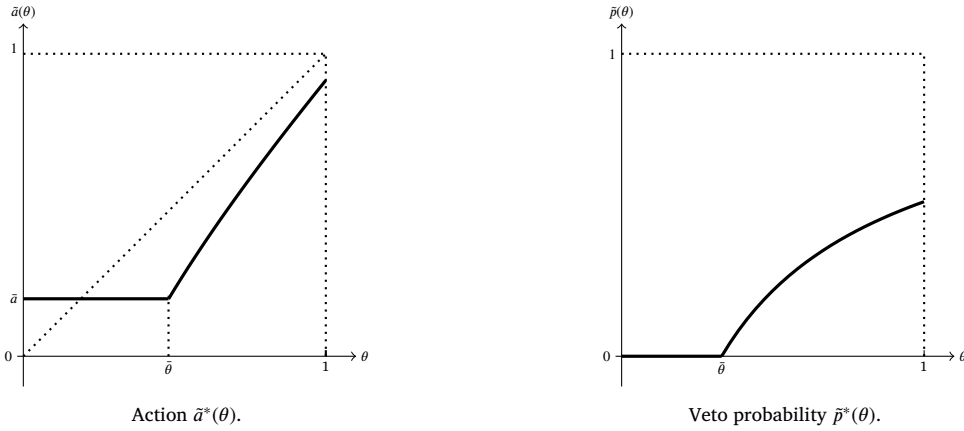


Fig. 1. Optimal veto mechanism.

Corollary 1. Among all valuable optimal veto mechanisms,

- (a) $\tilde{a}^*(\theta)$ and $\tilde{p}^*(\theta)$ are weakly increasing in θ ;
- (b) $\tilde{a}^*(\theta) < \theta$ for all $\theta > \bar{\theta}$.

Proof. See Appendix A.2. $\square$

The intuition of Corollary 1(a) is as follows. As the state θ increases, the principal's preferred option gets higher and thus $\tilde{a}^*(\theta)$ is weakly increasing. On the other hand, to maintain the indifference constraints (I), the principal has to veto those higher options with higher probabilities. So $\tilde{p}^*(\cdot)$ is weakly increasing as well. Corollary 1(b) claims that in the optimal veto mechanism the proposed option is always lower than the state, and its intuition is as follows. First let $\tilde{a}(\theta) = \theta$ for $\theta > \bar{\theta}$ and consider a marginal decrease of the proposed option. The principal's utility conditional on not using veto decreases, but that loss is of second-order.⁸ On the other hand, constraints (I) imply that a lower proposal necessitates a lower veto probability. As the status quo option yields payoffs $u_0 < 0$ for the principal, the gain from a marginal downward deviation from the state-matching action is first-order. To sum up, it benefits the principal to design the options $\tilde{a}(\theta)$ lower than θ for $\theta > \bar{\theta}$.

3.2. When are veto mechanisms valuable?

We say a veto mechanism is *valuable* if it yields a higher payoff for the principal than that of delegating $\{\mathbb{E}_\mu(\theta)\}$. By Proposition 2, whether there exists a valuable veto mechanism depends on players' payoffs from the status quo option, u_0 and v_0 , and the mean of prior μ .

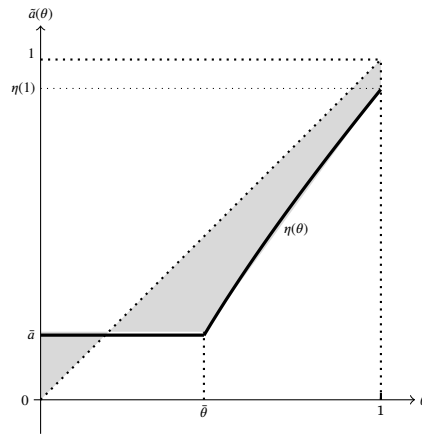
Effects of u_0 As $v_0 \leq 0$, the status quo option a_0 serves as a “punishment device.” That is, the principal uses the status quo option to balance the incentives of the agent to select higher options. Proposition 2 implies that the optimal veto mechanism is valuable if and only if $u_0 > (\mathbb{E}_\mu(\theta) - v_0)^2 - (1 - v_0)^2$. Put in other words, the status quo option serves as a valid punishment option only if it does not harm the principal that much. Indeed, when $u_0 \leq 2v_0 - 1$, there does not exist a valuable veto mechanism no matter what the principal's prior belief is. On the other hand, the restriction on u_0 for the existence of valuable veto mechanisms is not severe. It can be verified that a_0 can serve as a valid punishment option even when u_0 is significantly less than the principal's payoff from choosing his ex-ante preferred option $\mathbb{E}_\mu(\theta)$.⁹

Effects of v_0 The lower bound of u_0 for the existence of a valuable veto mechanism, $(\mathbb{E}_\mu(\theta) - v_0)^2 - (1 - v_0)^2$, is increasing in v_0 . The intuition is that when v_0 gets larger, the degree of punishment by veto becomes less severe for the agent. Then the principal has to punish the agent with higher probabilities in order to maintain the constraints (I). Therefore, the principal obtains the status quo payoff more often, making the veto mechanism harder to sustain.

Effects of prior μ The optimal veto mechanism is valuable only if u_0 is above the threshold $-(1 - v_0)^2$. In that case, the necessary and sufficient condition for the existence of valuable veto mechanisms can be written more parsimoniously as $\mathbb{E}_\mu(\theta) < \eta(1) \equiv$

⁸ As the principal's utility function is quadratic for $a \in [-M, M]$, the first-order derivative with respect to a at $a = \theta$ is zero: $\frac{\partial u(a, \theta)}{\partial a} \Big|_{a=\theta} = 0$ for all $\theta > \bar{\theta}$.

⁹ For an illustration, let the prior belief be uniform over Θ and set $v_0 = 0$. Then the status quo option can serve as a valid punishment option as long as u_0 is higher than $-3/4$, which is less than $-1/12$, the principal's payoff from choosing $\mathbb{E}_\mu(\theta)$.

Fig. 2. Pinning down the thresholds $\bar{a}$ and $\bar{\theta}$.

$\sqrt{(1 - v_0)^2 + u_0} + v_0$. Fixing u_0 and v_0 (and hence $\eta(1)$), the existence of valuable veto mechanisms depends only on the mean of the prior μ . As the principal's ex-post preferred option is $a = \theta$ and the agent always prefers higher actions regardless of the state, a higher $\mathbb{E}_\mu(\theta)$ indicates that the two players' interests are more aligned ex ante. In this sense, Proposition 2 implies that veto mechanisms are not valuable when the interests of two players are sufficiently aligned (i.e., $\mathbb{E}_\mu(\theta) > \eta(1)$). This finding is in contrast with the *Ally Principle* that greater alignment leads to more discretion in the delegation set.¹⁰ In our scenario, the principal leaves no discretion for the agent when players' preferences are sufficiently aligned, but leaves some discretion when preferences are sufficiently misaligned.

In a related paper, Kartik et al. (2021) document the invalidity of Ally Principle in a bargaining environment—a proposer (principal) delegates a set of lotteries for a vetoer (agent) to choose from while the vetoer can always choose to maintain the status quo besides the delegated options. In Kartik et al. (2021), the opposite of the Ally Principle is true: greater ex-ante alignment tends to make the principal leave less discretion for the agent (i.e., the delegation set gets strictly smaller). In a standard optimal delegation setting, Alonso and Matouschek (2008, Section 6.4) illustrate that the Ally Principle may fail when the optimal deterministic mechanism is not interval delegation. In this paper, the Ally Principle in general does not hold—the delegation set gets neither strictly smaller nor bigger as players have greater ex-ante preference alignment.¹¹

A deeper understanding of the effects of μ can be obtained by closely examining the threshold $\bar{a}$. The first-order condition determining $\bar{a}$ permits a graphical expression as in Fig. 2, where $\bar{a}$ and $\bar{\theta} \equiv \eta^{-1}(\bar{a})$ are the exact values that equalize the areas of the two shaded regions weighted by distribution μ . Since $\eta(1) < 1$, when the prior distribution μ puts too much mass on the interval $(1 - \epsilon, 1)$ for a sufficiently small $\epsilon > 0$, $\bar{a}$ would be above $\eta(1)$ and then the principal delegates $\{\mathbb{E}_\mu(\theta)\}$. Therefore, when $\mathbb{E}_\mu(\theta)$ is sufficiently high, there will be no valuable veto mechanisms.

3.3. Comparative statics

We derive two comparative statics to demonstrate how changes in u_0 and v_0 affect the optimal veto mechanism, focusing exclusively on valuable optimal veto mechanisms.

Proposition 3. Among valuable optimal veto mechanisms, if the value of status quo option for the principal u_0 is lower, then

- (a) $\eta(\theta)$ decreases for each possible θ ;
- (b) $\bar{p}^*(\theta)$ decreases whenever $\bar{p}^*(\theta) > 0$;
- (c) both $\bar{a}$ and $\bar{\theta}$ increase.

Proof. See Appendix A.3. $\square$

Fig. 3 illustrates how $\bar{a}^*(\theta)$ and $\bar{p}^*(\theta)$ change as principal's payoffs from the status quo option vary, with the dashed curves representing the case for a higher u_0 . The intuition for Proposition 3 is as follows. With a lower u_0 , the principal is less willing to maintain the status quo, leading to a decrease in the veto probability $\bar{p}^*(\theta)$ and an increase in the cutoff $\bar{\theta}$. To maintain the indifference constraints (I), the principal lowers the actions $\eta(\theta)$. The combination of a higher $\bar{\theta}$ and lower $\eta(\theta)$ leads to an increase of $\bar{a}$.

¹⁰ The Ally Principle has been shown to hold in a number of models in the political economy literature (See Bendor et al. (2001) and the discussions therein). Holmstrom (1980) demonstrates that the Ally Principle holds under general conditions when restricting to interval delegations.

¹¹ It is worth noting that in most papers where the Ally Principle holds, the degree of preferences misalignment is measured by some bias parameter and is constant across different states. In contrast, both Kartik et al. (2021) and our paper assume one player having state-independent preferences, and thus the degree of (ex-ante) preferences alignment is belief-dependent. In the counter-example of Alonso and Matouschek (2008), players' preference misalignment is not constant across states and is larger at higher states.

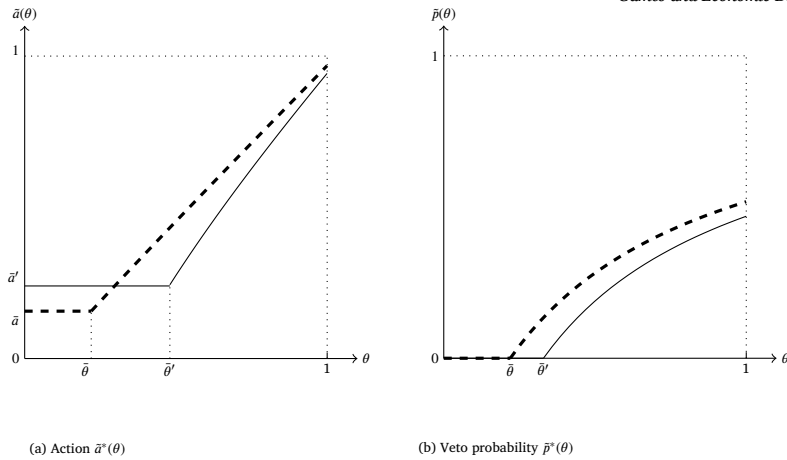


Fig. 3. Changes of $\bar{a}^*(\theta)$ and $\bar{p}^*(\theta)$ when u_0 varies.

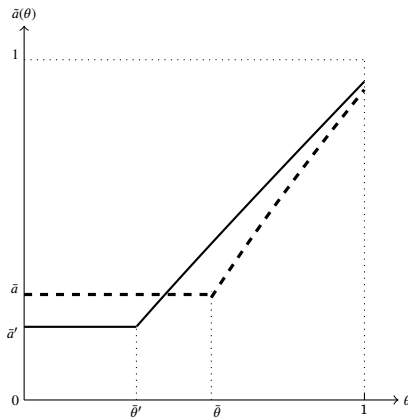


Fig. 4. Changes of $\bar{a}^*(\theta)$ when v_0 varies.

In Proposition 4, we derive comparative statics concerning the value of the status quo option for the agent. Unlike the case with varying u_0 , changes of the optimal veto probability function $\bar{p}^*(\theta)$ may or may not be uniform as v_0 varies.¹²

Proposition 4. Among valuable optimal veto mechanisms, if the value of status quo option for the agent v_0 is lower, then

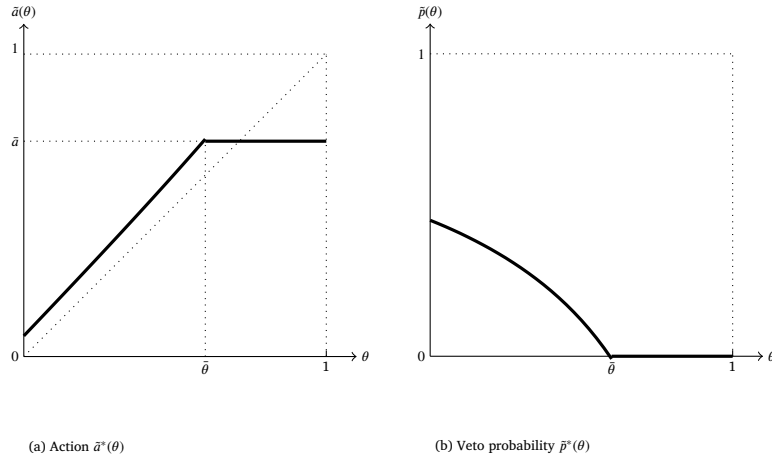
- (a) $\eta(\theta)$ increases for each possible θ ;
- (b) both $\bar{a}$ and $\bar{\theta}$ decrease.

Proof. See Appendix A.4. $\square$

Fig. 4 illustrates how $\bar{a}^*(\theta)$ changes as the agent's payoffs from the status quo option vary, with the dashed curve representing the optimal action rule for a higher v_0 . Intuitively, lower v_0 tends to lead to lower agent equilibrium payoff $\bar{a}$ as the agent's payoffs from the status quo option get lower. Meanwhile, lower v_0 implies that punishment by veto becomes more severe for the agent. To maintain the indifference constraints (T), the principal tends to increase the separating part of optimal actions $\eta(\theta)$. Since $\bar{a}$ decreases and $\eta(\theta)$ increases, the cutoff state $\bar{\theta} \equiv \eta^{-1}(\bar{a})$ decreases.

Lastly, we briefly discuss the comparative statics regarding the prior belief. One might conjecture that, among valuable optimal veto mechanisms, the state cutoff $\bar{\theta}$ changes monotonically when usual stochastic orderings (such as first-order stochastic dominance and likelihood ratios) are imposed on the prior beliefs. Nevertheless, this does not hold in our setting. In Appendix A.6, we provide numeric examples demonstrating that when μ_1 dominates μ_2 in the sense of likelihood ratio (and hence first-order stochastic dominance), the state cutoff induced by prior μ_1 can be either higher or lower than that induced by prior μ_2 .

¹² This is illustrated with numerical examples in Appendix A.5.

Fig. 5. Optimal veto mechanism when $v_0 \geq 1$.

4. Discussions

In the previous analysis, we have assumed $v_0 \leq 0$. In this section, we discuss how a positive v_0 will affect the optimal mechanism. When $v_0 \geq 1$, the optimal mechanism resembles that of our main model, with the key difference being that a_0 acts as a reward rather than a punishment option. However, when $0 < v_0 < 1$, the structure of the optimal mechanism may change significantly.

4.1. Scenario 1: $v_0 \geq 1$

Suppose $v_0 \geq 1$. Then the status quo option serves as a reward option rather than a punishment option. As Proposition 1 does not depend on the value of v_0 , we can still focus on the class of veto mechanisms and derive for the optimal mechanism by solving sequential maximization problems $(\mathcal{M}_1)$ and $(\mathcal{M}_2)$ as in the main model. The only difference is that now the agent prefers the status quo option and v_0 becomes the upper bound for the agent's equilibrium payoff, whereas in the previous analysis v_0 is the lower bound. The following proposition summarizes the optimal veto mechanism in this case.

Proposition 5. Assume $v_0 \geq 1$.

(a) If $(v_0 - \mathbb{E}_\mu(\theta))^2 < v_0^2 + u_0$, define $\varphi(\theta) \equiv v_0 - \sqrt{(v_0 - \theta)^2 + u_0}$ for $\theta \leq v_0 - \sqrt{-u_0}$ and the optimal veto mechanism $(\tilde{a}^*(\theta), \tilde{p}^*(\theta))$ is

$$\tilde{a}^*(\theta) = \begin{cases} \varphi(\theta), & \text{if } \theta < \bar{\theta} \\ \bar{a}, & \text{if } \theta \geq \bar{\theta} \end{cases} \quad \tilde{p}^*(\theta) = \begin{cases} \frac{\bar{a} - \varphi(\theta)}{v_0 - \varphi(\theta)}, & \text{if } \theta < \bar{\theta} \\ 0, & \text{if } \theta \geq \bar{\theta} \end{cases}$$

where $\bar{\theta} = \varphi^{-1}(\bar{a})$ and $\bar{a}$ is determined by $\int_0^{\varphi^{-1}(\bar{a})} (\varphi(\theta) - \bar{a}) d\mu + \bar{a} - \mathbb{E}_\mu(\theta) = 0$. Moreover, $\bar{a} \in (\mathbb{E}_\mu(\theta), 1)$.

(b) If $(v_0 - \mathbb{E}_\mu(\theta))^2 \geq v_0^2 + u_0$, then the optimal veto mechanism is $\tilde{a}^*(\theta) = \mathbb{E}_\mu(\theta)$ and $\tilde{p}^*(\theta) = 0$; that is, the principal chooses the action $\mathbb{E}_\mu(\theta)$ regardless of the state.

Proof. See Appendix A.7. $\square$

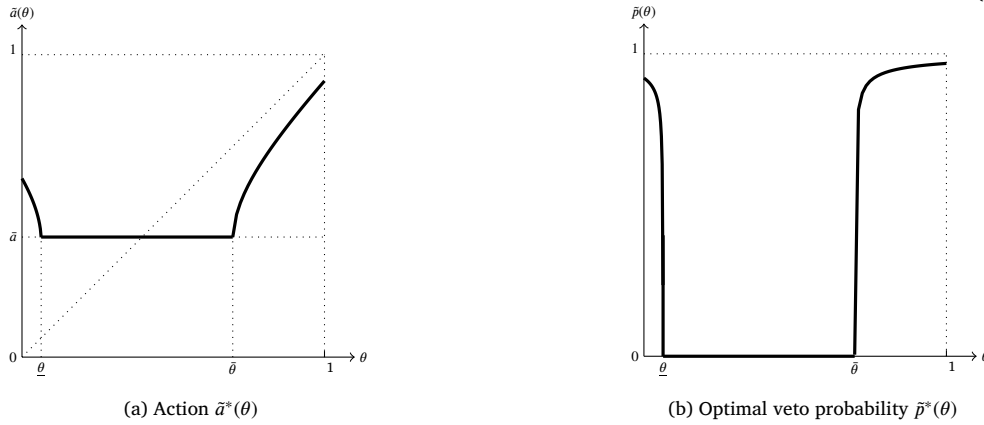
Fig. 5 illustrates the valuable optimal veto mechanism when $v_0 \geq 1$: the principal uses veto only when the state is below the threshold $\bar{\theta}$, and those corresponding actions $\tilde{a}^*(\theta)$ are higher than θ . Also, as opposed to the case of $v_0 \leq 0$, there exists no valuable veto mechanism when $\mathbb{E}_\mu(\theta)$ is sufficiently low (i.e., when players' preferences are sufficiently misaligned ex ante).

4.2. Scenario 2: $v_0 \in (0, 1)$

When $v_0 \in (0, 1)$, solving for the optimal mechanism generally becomes more complicated. Here, we characterize the optimal veto mechanism in two specific cases, where the solutions differ significantly from those in Propositions 2 and 5.

Let $v_0 = 0.38$, $u_0 = -0.1$ and the prior belief be the uniform distribution over Θ . The optimal action function $\tilde{a}^*(\theta)$ and the optimal veto rule $\tilde{p}^*(\theta)$ are U-shaped¹³:

¹³ Substituting the values of v_0 and u_0 into the expression of $\eta(\theta)$ as defined in Proposition 2 yields $\eta(\theta) = \sqrt{(\theta - 0.38)^2 - 0.1} + 0.38$ in this case.

Fig. 6. U-shaped $\tilde{a}^*(\theta)$ and $\tilde{p}^*(\theta)$.

$$\tilde{a}^*(\theta) = \begin{cases} \eta(\theta) & \text{if } \theta \in [0, \underline{\theta}); \\ \bar{a} & \text{if } \theta \in [\underline{\theta}, \bar{\theta}); \\ \eta(\theta) & \text{if } \theta \in [\bar{\theta}, 1], \end{cases} \quad \tilde{p}^*(\theta) = \begin{cases} \frac{\eta(\theta) - \bar{a}}{\eta(\theta) - v_0} & \text{if } \theta \in [0, \underline{\theta}); \\ 0 & \text{if } \theta \in [\underline{\theta}, \bar{\theta}); \\ \frac{\eta(\theta) - \bar{a}}{\eta(\theta) - v_0} & \text{if } \theta \in [\bar{\theta}, 1], \end{cases}$$

where the approximate values are given by $\underline{\theta} \approx 0.063$, $\bar{\theta} \approx 0.697$ and $\bar{a} \approx 0.397$. To obtain this result, we first divide the problem into two categories, $\hat{v} \geq v_0$ and $\hat{v} \leq v_0$, where $\hat{v}$ is the agent's expected payoff. For each case, we employ the two-step method to solve $(\mathcal{M}_1)$ and $(\mathcal{M}_2)$ sequentially. We find that it is optimal for the principal to set $\hat{v} \geq v_0$, and thus the status quo option serves as a stick rather than a carrot. Detailed derivations are relegated to Appendix A.8.1.

Fig. 6 illustrates the optimal veto mechanism: both $\tilde{a}^*(\theta)$ and $\tilde{p}^*(\theta)$ are U-shaped and the principal uses veto at both the higher and the lower states. Note that when $\theta < \underline{\theta}$, the principal's payoffs from those actions $\tilde{a}^*(\theta)$ are lower than that from the pooling action $\bar{a}$, yet the principal finds it optimal to take these higher actions. The reason for this phenomenon is as follows. When $\theta \in [0, \underline{\theta}]$, principal's payoff from the pooling action $\bar{a}$ is lower than u_0 . It follows that the principal has incentives to put more probability weight on the status quo option. As a result, the principal must also set $\tilde{a}(\theta)$ higher than $\bar{a}$ to maintain the indifference constraints (I).

As a final illustration, let $v_0 = \frac{1}{2}$, $u_0 > -\frac{1}{4}$ and the prior belief be the uniform distribution over Θ . In this case, the optimal mechanism is deterministic:

$$(\tilde{a}^*(\theta), \tilde{p}^*(\theta)) = \begin{cases} (v_0, 0) & \text{when } \theta \in [v_0 - \sqrt{-u_0}, v_0 + \sqrt{-u_0}) \\ (a', 1) & \text{otherwise} \end{cases}$$

where a' can be any non-status quo option distinct from v_0 . In other words, the principal always vetoes unless the proposed option is $a = v_0$. Detailed derivations are relegated to Appendix A.8.2.

5. Concluding remarks

This paper is motivated by the observation that both private and public organizations often use veto delegation to align incentives between principals and agents. We study an optimal delegation model in which it is optimal for the principal to use veto to elicit information from the agent. Our findings have implications for corporate governance and legislative studies, where veto delegation is prevalent.

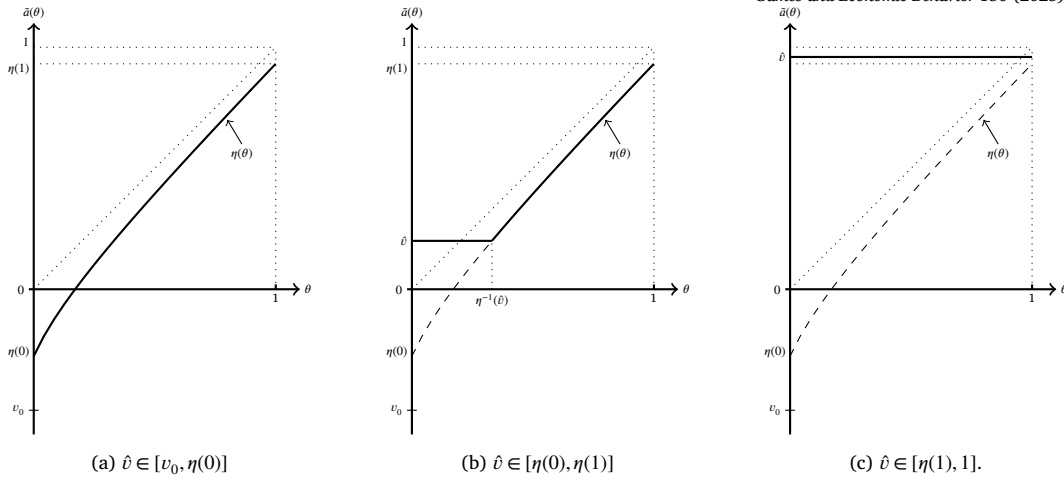
The optimality of veto mechanisms hinges on two key assumptions: state-independent agent preferences and the principal being more risk-averse than the agent towards non-status quo options. Further investigation into how the optimality of veto mechanisms depends on these conditions could be valuable. Additionally, we have made simplifying assumptions about players' payoff functions when characterizing the optimal mechanism. Relaxing these assumptions and exploring the optimal mechanism under different conditions could prove worthwhile for future research.

Declaration of competing interest

The authors Xiaoxiao Hu and Haoran Lei declare that they have no relevant or material financial interests that relate to the research described in this paper.

Appendix A

The Appendix contains the proofs and derivations which have been omitted from the main text.

Fig. 7. Principal's $\tilde{a}^*(\theta)$ for different $\hat{v}$ when $u_0 + v_0^2 \geq 0$.

A.1. Proof of Proposition 2

We first solve the maximization problem $(\mathcal{M}_1)$ at state θ for some fixed agent utility $\hat{v}$. Constraints (I) yield:

$$\tilde{p}(\theta) = \frac{\tilde{a}(\theta) - \hat{v}}{\tilde{a}(\theta) - v_0}. \quad (2)$$

Equation (2) implies that $\tilde{p}(\theta) \leq 1$ is equivalent to $\hat{v} \geq v_0$ and that $\tilde{p}(\theta) \geq 0$ is equivalent to $\hat{v} \leq \tilde{a}(\theta)$. Substitute Equation (2) into principal's utility function, and problem $(\mathcal{M}_1)$ is reduced to:

$$\max_{\tilde{a}(\theta) \in [-M, M]} \left(\frac{\tilde{a}(\theta) - \hat{v}}{\tilde{a}(\theta) - v_0} \right) u_0 - \left(\frac{\hat{v} - v_0}{\tilde{a}(\theta) - v_0} \right) (\tilde{a}(\theta) - \theta)^2 \quad \text{subject to } \tilde{a}(\theta) \geq \hat{v}. \quad (3)$$

The solution to problem (3) is

$$\tilde{a}^*(\theta) = \begin{cases} \max\{\hat{v}, \eta(\theta)\} & \text{if } (\theta - v_0)^2 + u_0 \geq 0 \\ \hat{v} & \text{otherwise} \end{cases} \quad (4)$$

where $\eta(\theta) = \sqrt{(\theta - v_0)^2 + u_0} + v_0$.

Then, we solve for the optimal $\hat{v}$ at state θ given the action rule specified by Equation (4). Specifically, we derive the optimal $\hat{v}$ for different values of u_0 and v_0 in the following scenarios:

- I. $u_0 + v_0^2 \geq 0$. In this scenario, $u_0 + (\theta - v_0)^2 \geq 0$ for all $\theta \in \Theta$.
- II. $u_0 + v_0^2 < 0$ and $u_0 + (1 - v_0)^2 > 0$. In this scenario, $u_0 + (\theta - v_0)^2 < 0$ for $\theta \in [0, \sqrt{-u_0} + v_0)$ and $u_0 + (\theta - v_0)^2 \geq 0$ for $\theta \in [\sqrt{-u_0} + v_0, 1]$.
- III. $u_0 + (1 - v_0)^2 \leq 0$. In this scenario, $u_0 + (\theta - v_0)^2 \leq 0$ for all $\theta \in \Theta$.

Scenario I: $u_0 + v_0^2 \geq 0$

We first explicitly write out principal's optimal action $\tilde{a}^*(\theta)$ specified by Equation (4) for three different ranges of $\hat{v}$ (Fig. 7), and then find the optimal $\hat{v}$ within each range. Finally, we compare the solutions in different ranges and find the global maximum.

When $\hat{v} \in [v_0, \eta(0)]$, principal's optimal action is $\tilde{a}^*(\theta) = \eta(\theta)$, and the veto probability is $\tilde{p}^*(\theta) = \frac{\eta(\theta) - \hat{v}}{\eta(\theta) - v_0}$. Principal's optimization problem is reduced to

$$\max_{\hat{v} \geq 0} \Gamma_1(\hat{v}) \equiv \int_0^1 \left(\frac{\eta(\theta) - \hat{v}}{\eta(\theta) - v_0} \right) u_0 - \left(\frac{\hat{v} - v_0}{\eta(\theta) - v_0} \right) (\eta(\theta) - \theta)^2 d\mu.$$

Since $\Gamma'_1(\hat{v}) = 2 \int_0^1 (\theta - v_0) - \sqrt{(\theta - v_0)^2 + u_0} d\mu > 0$ for all $\hat{v} \in [v_0, \eta(0)]$, the objective function $\Gamma_1(\hat{v})$ attains its maximum at $\hat{v}^* = \eta(0)$.

When $\hat{v} \in [\eta(0), \eta(1)]$, principal's optimal action is

$$\tilde{a}^*(\theta) = \begin{cases} \hat{v}, & \text{if } \theta \leq \eta^{-1}(\hat{v}); \\ \eta(\theta), & \text{if } \theta > \eta^{-1}(\hat{v}), \end{cases}$$

and the veto probability is

$$\bar{p}^*(\theta) = \begin{cases} 0, & \text{if } \theta \leq \eta^{-1}(\hat{v}); \\ \frac{\eta(\theta) - \hat{v}}{\eta(\theta) - v_0}, & \text{if } \theta > \eta^{-1}(\hat{v}). \end{cases}$$

Principal's maximization problem is reduced to

$$\max_{\hat{v}} \Gamma_2(\hat{v}) \equiv \int_0^{\eta^{-1}(\hat{v})} -(\hat{v} - \theta)^2 d\mu + \int_{\eta^{-1}(\hat{v})}^1 \left[\left(\frac{\eta(\theta) - \hat{v}}{\eta(\theta) - v_0} \right) u_0 - \left(\frac{\hat{v} - v_0}{\eta(\theta) - v_0} \right) (\eta(\theta) - \theta)^2 \right] d\mu$$

The first-order derivative is $\Gamma'_2(\hat{v}) = -2 \left(\int_{\eta^{-1}(\hat{v})}^1 \eta(\theta) - \hat{v} d\mu + \hat{v} - \mathbb{E}_\mu(\theta) \right)$. To find the sign of $\Gamma'_2(\hat{v})$, first note that $\Gamma'_2(\hat{v})$ is decreasing: $\Gamma''_2(\hat{v}) = -2 \int_0^{\eta^{-1}(\hat{v})} 1 d\mu < 0$. At the two end points $\eta(0)$ and $\eta(1)$, we have $\Gamma'_2(\eta(0)) = -2 \int_0^1 \eta(\theta) - \theta d\mu > 0$ and $\Gamma'_2(\eta(1)) = -2(\eta(1) - \mathbb{E}_\mu(\theta))$. Then,

1. when $\eta(1) \leq \mathbb{E}_\mu(\theta)$, we have $\Gamma'_2(\eta(1)) \geq 0$ and the solution is $\hat{v}^* = \eta(1)$;
2. when $\eta(1) > \mathbb{E}_\mu(\theta)$, we have $\Gamma'_2(\eta(1)) < 0$ and the solution is $\hat{v}^* = \bar{v}$ where $\bar{v}$ is determined by $\Gamma'_2(\bar{v}) = 0$.

Furthermore,

$$\Gamma'_2(0) = -2 \left(\int_{\eta^{-1}(0)}^1 \eta(\theta) d\mu - \mathbb{E}_\mu(\theta) \right) = 2 \left(\int_{\eta^{-1}(0)}^1 \theta - \eta(\theta) d\mu + \int_0^{\eta^{-1}(0)} \theta d\mu \right) > 0,$$

$$\Gamma'_2(\mathbb{E}_\mu(\theta)) = -2 \left(\int_{\eta^{-1}(\mathbb{E}_\mu(\theta))}^1 \eta(\theta) - \mathbb{E}_\mu(\theta) d\mu \right) < 0.$$

Since $\Gamma''_2(\hat{v}) < 0$, we have $\bar{v} \in (0, \mathbb{E}_\mu(\theta))$.

Lastly, when $\hat{v} \in [\eta(1), 1]$, principal's optimal action is $\tilde{a}^*(\theta) = \hat{v}$ and the veto probability is $\bar{p}^*(\theta) = 0$. Principal's optimization problem is reduced to

$$\max_{\hat{v}} \Gamma_3(\hat{v}) \equiv - \int_0^1 (\hat{v} - \theta)^2 d\mu.$$

The first-order derivative is $\Gamma'_3(\hat{v}) = -2(\hat{v} - \mathbb{E}_\mu(\theta))$. Note that $\Gamma''_3(\hat{v}) = -2 < 0$ and at the two end points we have:

$$\Gamma'_3(1) = -2(1 - \mathbb{E}_\mu(\theta)) < 0, \quad \Gamma'_3(\eta(1)) = -2(\eta(1) - \mathbb{E}_\mu(\theta)).$$

Then,

1. when $\eta(1) \leq \mathbb{E}_\mu(\theta)$, we have $\Gamma'_3(\eta(1)) \geq 0$ and the solution is $\hat{v}^* = \mathbb{E}_\mu(\theta)$;
2. when $\eta(1) > \mathbb{E}_\mu(\theta)$, we have $\Gamma'_3(\eta(1)) < 0$ and the solution is $\hat{v}^* = \eta(1)$.

To conclude, when $\eta(1) \leq \mathbb{E}_\mu(\theta)$, the optimal agent utility levels within the three ranges are given by:

$$\hat{v}^* = \begin{cases} \eta(0) & \text{if } \hat{v} \in [v_0, \eta(0)] \\ \eta(1) & \text{if } \hat{v} \in [\eta(0), \eta(1)] \\ \mathbb{E}_\mu(\theta) & \text{if } \hat{v} \in [\eta(1), 1] \end{cases}$$

Otherwise,

$$\hat{v}^* = \begin{cases} \eta(0) & \text{if } \hat{v} \in [v_0, \eta(0)] \\ \bar{v} & \text{if } \hat{v} \in [\eta(0), \eta(1)] \\ \eta(1) & \text{if } \hat{v} \in [\eta(1), 1] \end{cases}$$

where $\bar{v}$ is the solution to $\Gamma'_2(\bar{v}) = 0$.

Since the ranges overlap at the corner, we directly compare the solutions in the three ranges. The solution to Scenario I is as follows:

1. If $\eta(1) \leq \mathbb{E}_\mu(\theta)$, we have $\hat{v}^* = \mathbb{E}_\mu(\theta)$ and the optimal action is $\tilde{a}^*(\theta) = \mathbb{E}_\mu(\theta)$ for all θ .

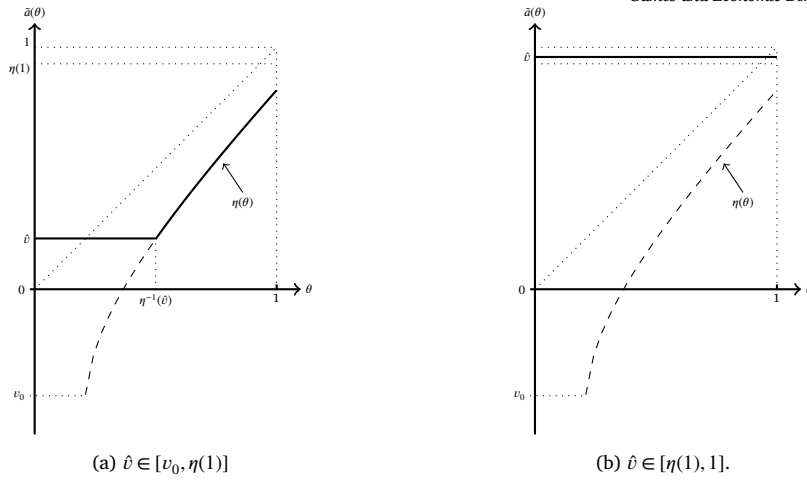


Fig. 8. Principal's $\tilde{a}^*(\theta)$ for different $\hat{v}$ when $u_0 + v_0^2 < 0$ and $u_0 + (1 - v_0)^2 > 0$.

2. If $\eta(1) > \mathbb{E}_\mu(\theta)$, we have $\hat{v}^* = \bar{v}$ where $\bar{v}$ is given by $\Gamma'_2(\bar{v}) = 0$ and the optimal action is as follows: when $\theta \in [0, \eta^{-1}(\bar{v})]$, $\tilde{a}^*(\theta) = \bar{v}$; when $\theta \in (\eta^{-1}(\bar{v}), 1]$, $\tilde{a}^*(\theta) = (\tilde{p}^*(\theta) \circ a_0, (1 - \tilde{p}^*(\theta)) \circ \eta(\theta))$ where $\tilde{p}^*(\theta) = \frac{\eta(\theta) - \bar{v}}{\eta(\theta) - v_0}$.

Scenario II: $u_0 + v_0^2 < 0$ and $u_0 + (1 - v_0)^2 > 0$

Fig. 8 illustrates the two possible cases of this scenario: $\hat{v} \in [v_0, \eta(1)]$ and $\hat{v} \in [\eta(1), 1]$. The derivation is similar to that of Scenario I (and therefore omitted), and the solution is the same as that of Scenario I.

Scenario III: $u_0 + (1 - v_0)^2 \leq 0$

In this scenario, principal's optimal action as described in Equation (4) is reduced to $\tilde{a}^*(\theta) = \hat{v}$ for all $\theta \in [0, 1]$, and the veto probability is $\tilde{p}^*(\theta) = 0$. Principal's maximization problem is reduced to

$$\max_{\hat{v}} - \int_0^1 (\hat{v} - \theta)^2 d\mu$$

The first-order condition yields:

$$-\int_0^1 2(\hat{v} - \theta) d\mu = 0 \implies \hat{v}^* = \mathbb{E}_\mu(\theta)$$

So, the optimal action is $\tilde{a}^*(\theta) = \mathbb{E}_\mu(\theta)$ for all θ .

Proposition 2 summarizes the three scenarios. Scenarios I and II are combined as they yield the same solution, with the overall prerequisite being $u_0 + (1 - v_0)^2 > 0$. Additionally, the condition $\eta(1) \leq \mathbb{E}_\mu(\theta)$ is equivalent to $u_0 + (1 - v_0)^2 \leq (\mathbb{E}_\mu(\theta) - v_0)^2$.

A.2. Proof of Corollary 1

Suppose the optimal veto mechanism is valuable. Then it has the semi-separating form as below:

1. When $\theta \in [0, \eta^{-1}(\bar{a})]$, $\tilde{a}^*(\theta) = \bar{a}$ and $\tilde{p}^*(\theta) = 0$.
2. When $\theta \in (\eta^{-1}(\bar{a}), 1]$, $\tilde{a}^*(\theta) = \eta(\theta) = \sqrt{(\theta - v_0)^2 + u_0} + v_0$. Taking derivative with respect to θ yields:

$$\eta'(\theta) = \frac{\theta - v_0}{\sqrt{(\theta - v_0)^2 + u_0}} > 0.$$

Taking derivative of $\tilde{p}^*(\theta) = \frac{\eta(\theta) - \bar{a}}{\eta(\theta) - v_0}$ with respect to θ gives

$$\tilde{p}^{*'}(\theta) = \frac{\eta'(\theta)(\eta(\theta) - v_0) - \eta'(\theta)(\eta(\theta) - \bar{a})}{(\eta(\theta) - v_0)^2} = \frac{\eta'(\theta)(\bar{a} - v_0)}{(\eta(\theta) - v_0)^2} > 0.$$

Therefore, $\tilde{a}^*(\theta)$ and $\tilde{p}^*(\theta)$ are constant over $\theta \in [0, \eta^{-1}(\bar{a})]$ and are strictly increasing over $\theta \in (\eta^{-1}(\bar{a}), 1]$.

As for part (b), fix some $\theta \in (\eta^{-1}(\bar{a}), 1]$. Since $u_0 < 0$, we have $\sqrt{(\theta - v_0)^2 + u_0} < \theta - v_0$. Therefore, $\eta(\theta) = \sqrt{(\theta - v_0)^2 + u_0} + v_0 < \theta$.

A.3. Proof of Proposition 3

Recall $\eta(\theta) = \sqrt{u_0 + (\theta - v_0)^2} + v_0$. Partial derivative of $\eta(\theta)$ with respect to u_0 gives

$$\frac{\partial \eta(\theta)}{\partial u_0} = \frac{1}{2\sqrt{u_0 + (\theta - v_0)^2}} > 0 \text{ for all } \theta \in [0, 1].$$

Therefore, $\eta(\theta)$ increases with u_0 for all $\theta \in [0, 1]$.

For $\bar{a}$, recall that $\bar{a}$ is obtained by:

$$\int_{\eta^{-1}(\bar{a})}^1 \eta(\theta) - \bar{a} d\mu + \bar{a} - \mathbb{E}(\theta) = 0.$$

Taking partial derivative with respect to u_0 yields:

$$\int_{\eta^{-1}(\bar{a})}^1 \frac{\partial \eta(\theta)}{\partial u_0} d\mu + \int_0^{\eta^{-1}(\bar{a})} \frac{\partial \bar{a}}{\partial u_0} d\mu = 0.$$

Since $\frac{\partial \eta(\theta)}{\partial u_0} > 0$ for all $\theta \in [0, 1]$, we obtain $\frac{\partial \bar{a}}{\partial u_0} < 0$.

For $\tilde{p}^*(\theta)$, recall that $\tilde{p}^*(\theta) = 1 - \frac{\bar{a} - v_0}{\eta(\theta) - v_0}$. Taking derivative with respect to u_0 gives

$$\frac{\partial \tilde{p}^*(\theta)}{\partial u_0} = - \left(\frac{\frac{\partial \bar{a}}{\partial u_0} (\eta(\theta) - v_0) - (\bar{a} - v_0) \frac{\partial \eta(\theta)}{\partial u_0}}{(\eta(\theta) - v_0)^2} \right).$$

Since $\frac{\partial \bar{a}}{\partial u_0} < 0$, $\eta(\theta) - v_0 > 0$, $\bar{a} - v_0 > 0$ and $\frac{\partial \eta(\theta)}{\partial u_0} > 0$, we have $\frac{\partial \tilde{p}^*(\theta)}{\partial u_0} > 0$.

A.4. Proof of Proposition 4

Recall $\eta(\theta) = \sqrt{u_0 + (\theta - v_0)^2} + v_0$. Partial derivative of $\eta(\theta)$ with respect to v_0 yields:

$$\frac{\partial \eta(\theta)}{\partial v_0} = - \left(\frac{\theta - v_0}{\sqrt{(\theta - v_0)^2 + u_0}} - 1 \right) < 0 \text{ for all } \theta \in [0, 1].$$

As for $\bar{a}$, recall that it is obtained by the FOC. Letting Γ'_2 equal to 0 yields:

$$\int_{\eta^{-1}(\bar{a})}^1 (\eta(\theta) - \bar{a}) d\mu + \bar{a} - \mathbb{E}(\theta) = 0.$$

Taking partial derivative with respect to v_0 yields:

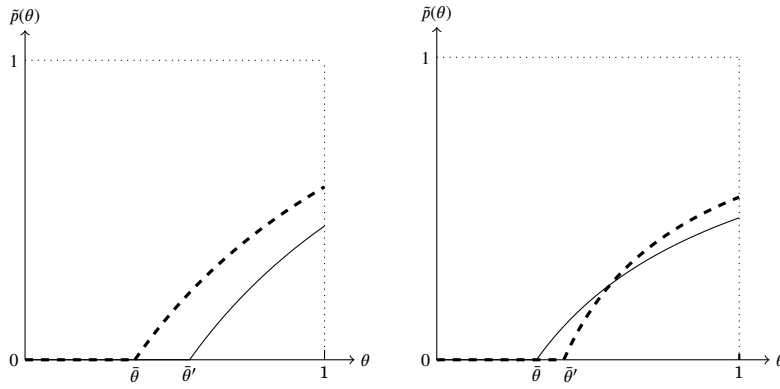
$$\int_{\eta^{-1}(\bar{a})}^1 \frac{\partial \eta(\theta)}{\partial v_0} d\mu + \int_0^{\eta^{-1}(\bar{a})} \frac{\partial \bar{a}}{\partial v_0} d\mu = 0.$$

As we have established that $\frac{\partial \eta(\theta)}{\partial v_0} < 0$ for all $\theta \in [0, 1]$, $\frac{\partial \bar{a}}{\partial v_0} > 0$ follows.

A.5. Veto probability $\tilde{p}^*(\theta)$ may not change uniformly as v_0 varies

We use two numeric examples to illustrate that $p^*(\theta)$ may or may not change uniformly when v_0 varies.

1. Suppose the density function is $g(\theta) = 1$ for $\theta \in \Theta$ and $u_0 = -0.9$. Consider the two cases: $v_0 = -0.9$ and $v'_0 = -0.3$. It follows from Proposition 2 that $p^*(\theta)$ weakly decreases for all $\theta \in \Theta$ as v_0 increases to v'_0 . The left panel of Fig. 9 illustrates this, where the dashed and solid curves represent the optimal veto probability under v_0 and v'_0 respectively.
2. Suppose the density function $g(\theta) = 1$ for $\theta \in \Theta$ and $u_0 = -0.2$. Consider the two cases: $v_0 = -0.5$ and $v'_0 = -0.3$. When the status quo option increases from v_0 to v'_0 , the optimal veto probability at state $\theta = 1$ increases from around 0.42 to around 0.47.

Fig. 9. Changes of $\bar{p}^*(\theta)$ when v_0 varies.

Meanwhile, the cutoff also increases from $\bar{\theta} \approx 0.44$ to $\bar{\theta}' \approx 0.49$. The right panel of Fig. 9 illustrates this, where the dashed and solid curves represent the optimal veto probability under v_0 and v'_0 respectively.

A.6. Insufficiency of the likelihood ratio conditions

We use numeric examples to illustrate that a higher likelihood ratio cannot guarantee either a higher or a lower $\bar{a}$. First, consider two distributions on $\Theta = [0, 1]$ with density functions g_L and g_H as follows:

$$g_L(\theta) = \begin{cases} 1/3 & \text{if } \theta \in [0, 0.9]; \\ 140(1 - \theta) & \text{if } \theta \in (0.9, 1]. \end{cases}$$

$$g_H(\theta) = \begin{cases} 1/3 & \text{if } \theta \in [0, 0.9]; \\ 140(\theta - 0.9) & \text{if } \theta \in (0.9, 1]. \end{cases}$$

It follows that g_H dominates g_L in the sense of likelihood ratio. Fixing $u_0 = -0.2$ and $v_0 = -0.6$, the corresponding threshold values are $\bar{\theta}_H \approx 0.647$ and $\bar{\theta}_L \approx 0.652$. So $\bar{\theta}_H < \bar{\theta}_L$.

On the other hand, consider two density functions $\hat{g}_L$ and $\hat{g}_H$ as follows:

$$\hat{g}_L(\theta) = 1 \text{ and } \hat{g}_H(\theta) = 10\theta^9, \forall \theta \in [0, 1].$$

It follows that $\hat{g}_H$ dominates $\hat{g}_L$ in the sense of likelihood ratio. Fixing $u_0 = -0.2$ and $v_0 = -0.6$, the corresponding threshold values are $\bar{\theta}'_H \approx 0.970$ and $\bar{\theta}'_L \approx 0.424$. So $\bar{\theta}'_H > \bar{\theta}'_L$.

A.7. Proof of Proposition 5

The proof of Proposition 5 is similar to that of Proposition 2.

Step 1: Pointwise Optimization

Fixing $\hat{v} \leq v_0$, constraints (I) imply

$$\bar{p}(\theta) = \frac{\hat{v} - \bar{a}(\theta)}{v_0 - \bar{a}(\theta)}. \quad (5)$$

Equation (5) implies that $\bar{p}(\theta) \leq 1$ is equivalent to $\hat{v} \leq \bar{a}(\theta)$ and that $\bar{p}(\theta) \geq 0$ is equivalent to $\hat{v} \geq \bar{a}(\theta)$. Substituting Equation (5) into the principal's utility, the optimization problem $(\mathcal{M}_1)$ is reduced to

$$\begin{aligned} \max_{\bar{a}(\theta) \in A} & \left(\frac{\hat{v} - \bar{a}(\theta)}{v_0 - \bar{a}(\theta)} \right) u - \left(\frac{v_0 - \hat{v}}{v_0 - \bar{a}(\theta)} \right) (\bar{a}(\theta) - \theta)^2 \\ \text{subject to } & \bar{a}(\theta) \leq \hat{v}, \forall \theta \in \Theta. \end{aligned} \quad (6)$$

The solution to problem (6) is

$$\bar{a}^*(\theta) = \begin{cases} \hat{v} & \text{if } (v_0 - \theta)^2 + u_0 < 0, \\ \min\{\hat{v}, \varphi(\theta)\} & \text{otherwise} \end{cases} \quad (7)$$

where $\varphi(\theta) = v_0 - \sqrt{(v_0 - \theta)^2 + u_0}$.

Step 2: Optimal $\hat{v}$

Based on Equation (7), we derive the optimal $\hat{v}$ for different values of u_0 and v_0 characterized by the following three cases:

- (i) $u_0 + (v_0 - 1)^2 \geq 0$. In this case, $u_0 + (v_0 - \theta)^2 \geq 0$ for all $\theta \in [0, 1]$.
- (ii) $u_0 + v_0^2 > 0$ and $u_0 + (v_0 - 1)^2 < 0$. In this case, $u_0 + (v_0 - \theta)^2 \geq 0$ for $\theta \in [0, \sqrt{-u_0} + v_0]$ and $u_0 + (v_0 - \theta)^2 < 0$ for $\theta \in (\sqrt{-u_0} + v_0, 1]$.
- (iii) $u_0 + v_0^2 \leq 0$. In this case, $u_0 + (v_0 - \theta)^2 \leq 0$ for all $\theta \in [0, 1]$.

The solutions of case (i) and case (ii) are the same:

1. If $\varphi(0) \geq \mathbb{E}_\mu(\theta)$, $\hat{v}^* = \mathbb{E}_\mu(\theta)$. The actions are pooled at $\tilde{a}(\theta) = \mathbb{E}_\mu(\theta)$.
2. If $\varphi(0) < \mathbb{E}_\mu(\theta)$, $\hat{v}^* = \bar{v}$ where $\bar{v}$ is given by

$$\int_0^{\varphi^{-1}(\bar{v})} \varphi(\theta) - \bar{v} d\mu + \bar{v} - \mathbb{E}_\mu(\theta) = 0$$

The action function $\tilde{a}(\theta)$ is as follows:

- when $\theta \in [0, \varphi^{-1}(\bar{v})]$, the default action a_0 is chosen with probability $\frac{\bar{v} - \varphi(\theta)}{v_0 - \varphi(\theta)}$ and the action $\tilde{a}(\theta) = \varphi(\theta)$ is chosen with the complementary probability;
- when $\theta \in [\varphi^{-1}(\bar{v}), 1]$, $\tilde{a}(\theta) = \bar{v}$.

We could further show that $\bar{v} \in (\mathbb{E}_\mu(\theta), 1)$. Let $f(\hat{v}) = \int_0^{\varphi^{-1}(\hat{v})} \varphi(\theta) - \hat{v} d\mu + \hat{v} - \mathbb{E}_\mu(\theta)$. Then, $\bar{v}$ is the solution to $f(\hat{v}) = 0$. $f(\hat{v})$ is strictly increasing:

$$f'(\hat{v}) = \int_{\varphi^{-1}(\hat{v})}^1 1 d\mu(\theta) > 0.$$

At the two points $\hat{v} = \mathbb{E}_\mu(\theta)$ and $\hat{v} = 1$:

$$\begin{aligned} f(\mathbb{E}_\mu(\theta)) &= - \int_0^{\varphi^{-1}(\mathbb{E}_\mu(\theta))} \mathbb{E}_\mu(\theta) - \varphi(\theta) d\mu < 0 \\ f(1) &= \int_0^{\varphi^{-1}(1)} \varphi(\theta) - 1 d\mu + 1 - \mathbb{E}_\mu(\theta) \\ &= \int_0^{\varphi^{-1}(1)} \varphi(\theta) - \theta d\mu + \int_{\varphi^{-1}(1)}^1 1 - \theta d\mu > 0 \end{aligned}$$

It follows that $\bar{v} \in (\mathbb{E}_\mu(\theta), 1)$.

The condition for cases (i) and (ii) combined is $u_0 + v_0^2 > 0$. Furthermore, the condition for the pooled action $\varphi(0) \geq \mathbb{E}_\mu(\theta)$ is equivalent to $(v_0 - \mathbb{E}_\mu(\theta))^2 \geq u_0 + v_0^2$. Therefore,

- When $(v_0 - \mathbb{E}_\mu(\theta))^2 \geq u_0 + v_0^2 > 0$, the actions are pooled at $\tilde{a}^*(\theta) = \mathbb{E}_\mu(\theta)$;
- When $(v_0 - \mathbb{E}_\mu(\theta))^2 < u_0 + v_0^2$, the optimal action rule $(\tilde{a}^*, \tilde{p}^*)$ is as characterized in part (a) of Proposition 5.

For case (iii), the principal optimally sets $\hat{v}^* = \mathbb{E}_\mu(\theta)$, resulting in pooled action $\tilde{a}^*(\theta) = \mathbb{E}_\mu(\theta)$.

Proposition 5 summarizes the analysis above: the cases $(v_0 - \mathbb{E}_\mu(\theta))^2 \geq u_0 + v_0^2 > 0$ and $u_0 + v_0^2 \leq 0$ both result in the pooled action and are stated in part (b) of the proposition; the case $(v_0 - \mathbb{E}_\mu(\theta))^2 < u_0 + v_0^2$ corresponds to part (a) of the proposition.

A.8. Derivations omitted in Section 4.2**A.8.1. U-shaped $\tilde{a}^*(\theta)$ and $\tilde{p}^*(\theta)$**

The parameters under concern are $v_0 = 0.38$, $u_0 = -0.1$ and $g(\theta) = 1$ for all $\theta \in [0, 1]$. Proposition 1 still applies. We can focus on the class of veto mechanisms when searching for a principal's optimal mechanism and decompose principal's maximization problem

into two sequential problems $(\mathcal{M}_1)$ and $(\mathcal{M}_2)$. The solution to $(\mathcal{M}_1)$ is different when $\hat{v} \geq v_0$ and $\hat{v} \leq v_0$. And we consider the two cases separately.

Case I: $\hat{v} \geq v_0$ When $\hat{v} \geq v_0$, we have $\tilde{a}(\theta) \geq v_0$ for all $\theta \in [0, 1]$. The solution to $(\mathcal{M}_1)$ is the same as the solution for the case $v_0 \leq 0$:

$$\tilde{a}^*(\theta) = \begin{cases} \hat{v} & \text{if } (\theta - v_0)^2 + u_0 < 0, \\ \max\{\hat{v}, \eta(\theta)\} & \text{otherwise} \end{cases}$$

where $\eta(\theta) = \sqrt{(\theta - v_0)^2 + u_0} + v_0$. Plugging in the parameter values, we have

1. If $\hat{v} \in [v_0, \sqrt{0.0444} + 0.38]$,

$$\tilde{a}^*(\theta) = \begin{cases} \hat{v} & \text{if } \underline{\theta} < \theta < \bar{\theta}, \\ \sqrt{(\theta - 0.38)^2 - 0.1} + 0.38 & \text{otherwise} \end{cases}$$

where $\underline{\theta} = 0.38 - \sqrt{(\hat{v} - 0.38)^2 + 0.1}$ and $\bar{\theta} = 0.38 + \sqrt{(\hat{v} - 0.38)^2 + 0.1}$.

2. If $\hat{v} \in [\sqrt{0.0444} + 0.38, \sqrt{0.2844} + 0.38]$,

$$\tilde{a}^*(\theta) = \begin{cases} \hat{v} & \text{if } \theta < \bar{\theta}, \\ \sqrt{(\theta - 0.38)^2 - 0.1} + 0.38 & \text{otherwise} \end{cases}$$

where $\bar{\theta} = 0.38 + \sqrt{(\hat{v} - 0.38)^2 + 0.1}$.

3. If $\hat{v} \geq \sqrt{0.2844} + 0.38$, $\tilde{a}^*(\theta) = \hat{v}$.

Next, we solve the maximization problem $(\mathcal{M}_2)$. The optimal $\hat{v}$ is within the range $\hat{v} \in [v_0, \sqrt{0.0444} + 0.38]$ and given by the below FOC:

$$\int_0^{\underline{\theta}} \eta(\theta) - \hat{v} d\theta + \int_{\bar{\theta}}^1 \eta(\theta) - \hat{v} d\theta + \hat{v} - \frac{1}{2} = 0 \quad (8)$$

where $\eta(\theta) = \sqrt{(\theta - 0.38)^2 - 0.1} + 0.38$, $\underline{\theta} = 0.38 - \sqrt{(\hat{v} - 0.38)^2 + 0.1}$ and $\bar{\theta} = 0.38 + \sqrt{(\hat{v} - 0.38)^2 + 0.1}$ as calculated when solving $(\mathcal{M}_1)$. Solving Equation (8) numerically gives $\hat{v}^* \approx 0.397$, $\underline{\theta} \approx 0.063$ and $\bar{\theta} \approx 0.697$.

Case II: $\hat{v} \leq v_0$ When $\hat{v} \leq v_0$, we have $\tilde{a}(\theta) \leq v_0$ for all $\theta \in [0, 1]$. The solution to $(\mathcal{M}_1)$ is the same as the solution for the case $v_0 \geq 1$.

$$\tilde{a}^*(\theta) = \begin{cases} \hat{v} & \text{if } (v_0 - \theta)^2 + u_0 < 0, \\ \min\{\hat{v}, \varphi(\theta)\} & \text{otherwise} \end{cases}$$

where $\varphi(\theta) = v_0 - \sqrt{(v_0 - \theta)^2 + u_0}$. Plugging in the parameter values, we have

1. If $\hat{v} \in [0.38 - \sqrt{0.0444}, v_0]$,

$$\tilde{a}^*(\theta) = \begin{cases} \hat{v} & \text{if } \underline{\theta} < \theta < \bar{\theta}, \\ 0.38 - \sqrt{(\theta - 0.38)^2 - 0.1} & \text{otherwise} \end{cases}$$

where $\underline{\theta} = 0.38 - \sqrt{(\hat{v} - 0.38)^2 + 0.1}$ and $\bar{\theta} = 0.38 + \sqrt{(\hat{v} - 0.38)^2 + 0.1}$.

2. If $\hat{v} \in [0.38 - \sqrt{0.2844}, 0.38 - \sqrt{0.0444}]$,

$$\tilde{a}^*(\theta) = \begin{cases} \hat{v} & \text{if } \theta < \bar{\theta}, \\ 0.38 - \sqrt{(\theta - 0.38)^2 - 0.1} & \text{otherwise} \end{cases}$$

where $\bar{\theta} = 0.38 + \sqrt{(\hat{v} - 0.38)^2 + 0.1}$.

3. If $\hat{v} \leq 0.38 - \sqrt{0.2844}$, $\tilde{a}^*(\theta) = \hat{v}$.

Next, we solve the maximization problem $(\mathcal{M}_2)$. Note that the optimal $\hat{v}$ is within the range $\hat{v} \in [0.38 - \sqrt{0.0444}, v_0]$ and that the first-order derivative is positive for all $\hat{v} \in [0.38 - \sqrt{0.0444}, v_0 = 0.38]$:

$$2 \left[\int_0^{\underline{\theta}} \hat{v} - \varphi(\theta) d\theta + \int_{\bar{\theta}}^1 \hat{v} - \varphi(\theta) d\theta + \left(\frac{1}{2} - \hat{v} \right) \right] > 0.$$

Therefore, we obtain the corner solution $\hat{v}^* = v_0 = 0.38$.

Combining the two cases, the optimal solution is obtained when $\hat{v}^* \approx 0.397$.

A.8.2. Deterministic optimal mechanism

We follow the same procedures as in the previous proof of Appendix A.8.1.

Case I: $\hat{v} \geq v_0$ The solution to $(\mathcal{M}_1)$ is the same as the solution for the case $v_0 \leq 0$.

1. If $\hat{v} \in [v_0 = 0.5, \sqrt{0.25 + u_0} + 0.5]$,

$$\tilde{a}^*(\theta) = \begin{cases} \hat{v} & \text{if } \underline{\theta} \leq \theta \leq \bar{\theta}, \\ \sqrt{(\theta - 0.5)^2 + u_0} + 0.5 & \text{otherwise} \end{cases}$$

where $\underline{\theta} = 0.5 - \sqrt{(\hat{v} - 0.5)^2 - u_0}$ and $\bar{\theta} = 0.5 + \sqrt{(\hat{v} - 0.5)^2 - u_0}$.

2. If $\hat{v} \geq \sqrt{0.25 + u_0} + 0.5$, $\tilde{a}^*(\theta) = \hat{v}$.

Next, we solve the maximization problem $(\mathcal{M}_2)$. The optimal $\hat{v}$ is within the range $[v_0, \sqrt{0.25 + u_0} + 0.5]$ and the first-order derivative is negative for all $\hat{v} \in [v_0, \sqrt{0.25 + u_0} + 0.5]$:

$$-2 \left[\int_0^{\underline{\theta}} \eta(\theta) - \hat{v} d\theta + \int_{\bar{\theta}}^1 \eta(\theta) - \hat{v} d\theta + \left(\hat{v} - \frac{1}{2} \right) \right] < 0.$$

Therefore, we obtain a corner solution $\hat{v}^* = v_0 = 0.5$.

Case II: $\hat{v} \leq v_0$ The solution to $(\mathcal{M}_1)$ is the same as the solution for the case $v_0 \geq 1$.

1. If $\hat{v} \in [0.5 - \sqrt{0.25 + u_0}, v_0 = 0.5]$,

$$\tilde{a}^*(\theta) = \begin{cases} \hat{v} & \text{if } \underline{\theta} \leq \theta \leq \bar{\theta}, \\ 0.5 - \sqrt{(\theta - 0.5)^2 + u_0} & \text{otherwise} \end{cases}$$

where $\underline{\theta} = 0.5 - \sqrt{(\hat{v} - 0.5)^2 - u_0}$ and $\bar{\theta} = 0.5 + \sqrt{(\hat{v} - 0.5)^2 - u_0}$.

2. If $\hat{v} \leq 0.5 - \sqrt{0.25 + u_0}$, $\tilde{a}^*(\theta) = \hat{v}$.

Next, we solve the maximization problem $(\mathcal{M}_2)$. The optimal $\hat{v}$ is within the range $\hat{v} \in [0.5 - \sqrt{0.25 + u_0}, v_0]$ and the first-order derivative is positive for all $\hat{v} \in [0.5 - \sqrt{0.25 + u_0}, v_0]$:

$$2 \left[\int_0^{\underline{\theta}} \hat{v} - \varphi(\theta) d\theta + \int_{\bar{\theta}}^1 \hat{v} - \varphi(\theta) d\theta + \left(\frac{1}{2} - \hat{v} \right) \right] > 0.$$

Therefore, we obtain a corner solution $\hat{v}^* = v_0$.

Combining the two cases, we obtain that in the solution $\hat{v}^* = v_0$. Then $\underline{\theta}^* = 0.5 - \sqrt{-u_0}$ and $\bar{\theta}^* = 0.5 + \sqrt{-u_0}$. And both cases correspond to the veto probability:

$$\tilde{p}^*(\theta) = \begin{cases} 0 & \text{if } \underline{\theta}^* \leq \theta \leq \bar{\theta}^*, \\ 1 & \text{otherwise.} \end{cases}$$

Since the principal vetoes with probability 1 when $\theta \in [0, \underline{\theta}^*]$ and $\theta \in (\bar{\theta}^*, 1]$, $\tilde{a}^*(\theta)$ can take any value. When $\theta \in [\underline{\theta}^*, \bar{\theta}^*]$, we have $\tilde{a}^*(\theta) = \hat{v}^* = v_0$.

Data availability

No data was used for the research described in the article.

References

- Aghion, P., Tirole, J., 1997. Formal and real authority in organizations. *J. Polit. Econ.* 105 (1), 1–29.
 Alonso, R., Matouschek, N., 2008. Optimal delegation. *Rev. Econ. Stud.* 75 (1), 259–293.
 Amador, M., Bagwell, K., 2013. The theory of optimal delegation with an application to tariff caps. *Econometrica* 81 (4), 1541–1599.
 Baron, D.P., Myerson, R.B., 1982. Regulating a monopolist with unknown costs. *Econometrica* 50 (4), 911–930.

- Bartling, B., Fischbacher, U., 2012. Shifting the blame: on delegation and responsibility. *Rev. Econ. Stud.* 79 (1), 67–87.
- Bendor, J., Glazer, A., Hammond, T., 2001. Theories of delegation. *Annu. Rev. Pol. Sci.* 4 (1), 235–269.
- Chakraborty, A., Harbaugh, R., 2010. Persuasion by cheap talk. *Am. Econ. Rev.* 100 (5), 2361–2382.
- Chen, Y., 2022. Dynamic delegation with a persistent state. *Theor. Econ.* 17, 1589–1618.
- Dessein, W., 2002. Authority and communication in organizations. *Rev. Econ. Stud.* 69 (4), 811–838.
- Diehl, C., Kuzmics, C., 2021. The (non-)robustness of influential cheap talk equilibria when the sender's preferences are state independent. *Int. J. Game Theory*, 1–15.
- Doran, M., 2010. The closed rule. *Emory Law J.* 59, 1363.
- Frankel, A., 2016. Discounted quotas. *J. Econ. Theory* 166, 396–444.
- Gale, D., Hellwig, M., 1985. Incentive-compatible debt contracts: the one-period problem. *Rev. Econ. Stud.* 52 (4), 647–663.
- Glazer, J., Rubinstein, A., 2004. On optimal rules of persuasion. *Econometrica* 72 (6), 1715–1736.
- Glazer, J., Rubinstein, A., 2006. A study in the pragmatics of persuasion: a game theoretical approach. *Theor. Econ.* 1, 395–410.
- Hart, S., Kremer, L., Perry, M., 2017. Evidence games: truth and commitment. *Am. Econ. Rev.* 107 (3), 690–713.
- Holmstrom, Bengt, 1980. On the theory of delegation. *Discussion Papers*. <https://ideas.repec.org/p/nwu/cmsems/438.html>.
- Kahneman, D., Tversky, A., 1979. Prospect theory: an analysis of decision under risk. *Econometrica* 47 (2), 263–292.
- Kamenica, E., Gentzkow, M., 2011. Bayesian persuasion. *Am. Econ. Rev.* 101 (6), 2590–2615.
- Kartik, N., Kleiner, A., Van Weelden, R., 2021. Delegation in veto bargaining. *Am. Econ. Rev.* 111 (12), 4046–4087.
- Kolotilin, A., Zapechelnyuk, A., 2019. Persuasion meets delegation. *Working Paper*. <https://arxiv.org/abs/1902.02628>.
- Kováč, E., Mylovanov, T., 2009. Stochastic mechanisms in settings without monetary transfers: the regular case. *J. Econ. Theory* 144 (4), 1373–1395.
- Krishna, V., Morgan, J., 2001. Asymmetric information and legislative rules: some amendments. *Am. Polit. Sci. Rev.* 95 (2), 435–452.
- Lipnowski, E., Ravid, D., 2020. Cheap talk with transparent motives. *Econometrica* 88 (4), 1631–1660.
- Lubensky, D., Schmidbauer, E., 2018. Equilibrium informativeness in veto games. *Games Econ. Behav.* 109, 104–125.
- Marino, A.M., 2007. Delegation versus veto in organizational games of strategic communication. *J. Public Econ. Theory* 9 (6), 979–992.
- Martimort, D., Semenov, A., 2006. Continuity in mechanism design without transfers. *Econ. Lett.* 93 (2), 182–189.
- Melumad, N.D., Shibano, T., 1991. Communication in settings with no transfers. *Rand J. Econ.*, 173–198.
- Mylovanov, T., 2008. Veto-based delegation. *J. Econ. Theory* 138 (1), 297–307.