

Mylovanov, Tymofiy

Working Paper

Veto-Based Delegation

SFB/TR 15 Discussion Paper, No. 129

Provided in Cooperation with:

Free University of Berlin, Humboldt University of Berlin, University of Bonn, University of Mannheim, University of Munich, Collaborative Research Center Transregio 15: Governance and the Efficiency of Economic Systems

Suggested Citation: Mylovanov, Tymofiy (2005) : Veto-Based Delegation, SFB/TR 15 Discussion Paper, No. 129, Sonderforschungsbereich/Transregio 15 - Governance and the Efficiency of Economic Systems (GESY), München,
<https://doi.org/10.5282/ubm/epub.13422>

This Version is available at:

<https://hdl.handle.net/10419/94014>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.



GOVERNANCE AND THE EFFICIENCY
OF ECONOMIC SYSTEMS
GESY

Discussion Paper No. 129
Veto-Based Delegation
Tymofiy Mylovanov*

October 2004
January 2005

*Tymofiy Mylovanov, Department of Economics, University of Bonn, Adenauerallee 24-42, 53113 Bonn.
mylovanov@uni-bonn.de

Financial support from the Deutsche Forschungsgemeinschaft through SFB/TR 15 is gratefully acknowledged.

Sonderforschungsbereich/Transregio 15 · www.gesy.uni-mannheim.de
Universität Mannheim · Freie Universität Berlin · Humboldt-Universität zu Berlin · Ludwig-Maximilians-Universität München
Rheinische Friedrich-Wilhelms-Universität Bonn · Zentrum für Europäische Wirtschaftsforschung Mannheim

Speaker: Prof. Konrad Stahl, Ph.D. · Department of Economics · University of Mannheim · D-68131 Mannheim,
Phone: +49(0621)1812786 · Fax: +49(0621)1812785

Veto-Based Delegation.

Tymofiy Mylovanov *

First version: October 2004

This version: January 2005

Abstract

In a principal-agent model with hidden information and no monetary transfers, I establish the Veto-Power Principle: any incentive-compatible outcome can be implemented through veto-based delegation with an endogenously chosen default decision. This result demonstrates the exact nature of commitment powers required by the principal: (1) to design the default outcome and (2) to ensure that she has *almost no* formal control over the agent's decisions.

JEL codes: D78, D82, L22, M54.

Keywords: veto power, asymmetric information, principal-agent relationship, no monetary transfers.

1 Introduction

Decision making in real life frequently incorporates elements of veto-based delegation: an agent has the right to propose a decision, while a principal has the right to veto it. The US Congress, for example, has been using legislative veto, expedited congressional oversight, and other provisions that allow for disapproval of decisions delegated to executive agencies. The British parliament employs two procedures, affirmative and negative, for the review of delegated legislation. Under affirmative procedure a legislation has to be explicitly approved by the parliament, whereas under negative procedure a legislation becomes a law unless it is explicitly disapproved within 40 days. On some issues the European Union legislature (the Parliament and the Council) subject the decisions of the executive (the European Commission) to the

*Department of Economics, University of Bonn, Adenauerallee 24-42, 53113 Bonn. Email:mylovanov@uni-bonn.de. This paper is based on my dissertation, submitted to the Graduate School of the University of Wisconsin - Madison. I would like to thank my advisor Larry Samuelson for his continuous encouragement and support. I am also grateful to James Andreoni, Roy Gardner, Georg Gebhardt, Grigory Kosenok, Bart Lipman, Georg Noldeke, Lucia Quesada, Bill Sandholm, and Patrick Schmitz for numerous discussions and suggestions. Financial support from the project SFB 15, Projektbereich A is gratefully acknowledged.

approval of regulatory committees. In US open corporations the proposals by the board of directors about its new members, auditor choice, mergers, and new stock issues have to be approved by shareholders. In turn, the boards ratify major policy initiatives, and approve decisions on hiring, firing, and amount of compensation for important decision managers. The creditors of a corporation are often provided with guarantees that incurring other debt, making substantial capital acquisitions, increasing dividends, and changing the executive officers cannot be made without their consent.

These examples have several common attributes: (1) one party (an agent, he) has superior decision-relevant information and has the right and the responsibility to initiate decisions, while (2) the other party (a principal, she) has the right to reject the decision but typically does not have a formal right, is unqualified, constrained by time, or has to incur a high opportunity cost to develop a counter proposal. Furthermore, there are no payments between the parties that are directly tailored to the proposed decisions.

Although sometimes these elements of veto-based delegation are exogenously given (e.g., referendum, ratification of a treaty), they frequently appear to be a result of conscious design (e.g., during the Great Depression the legislative veto emerged as an innovative form of legislation tool purporting to overcome delays associated with the standard legislative process).¹ In such situations, the principal may be also capable of choosing a default decision, which is implemented in the absence of approval (e.g., Appropriation Acts of the US Congress usually specify the purposes and the amounts for these purposes that do not require further approval).²

In this paper I analyze the performance of veto-based delegation in the principal-agent environment with hidden information and no monetary transfers. I demonstrate that under some regularity condition any feasible outcome, and hence any desired outcome, can be implemented through veto-based delegation with a properly chosen default decision.

This result is a generalization of a simple observation. Consider a cheap-talk communication game (Crawford and Sobel [3]). In this game, any equilibrium has a partition structure and each element of this partition leads to a distinct choice. On the equilibrium path, the sender (the agent) chooses from among these different alternatives by sending a corresponding message; if the sender uses an out-of-equilibrium message, the receiver (the principal) takes a certain decision, say a , regardless of what the sender's private information might have been. One can rewrite this equilibrium as an outcome of veto-based delegation as follows: The agent proposes a decision. If

¹See Fisher [6] for discussion of the origins and history of the legislative veto.

²For instance, in Department of Homeland Security Appropriation Act 2006, H.R.2360, USD M 340 is appropriated for the development of Visitor and Immigrant Status Indicator Technology, of which USD M 160 may not be used until the Committees on Appropriations receive and approve a plan for expenditure, whereas *none* of USD M 456 appropriated for customs and border protection automated systems may be spent without approval of the Committees. Similarly, provisions determining what should and what should not be approved are contained in sections dealing with Air and Marine Operations, Customs Enforcement Automated Systems, Domestic Nuclear Detection Office, and General Provisions.

this decision could have been induced through some message in the original game, the principal approves it. Otherwise, she vetoes it using action a .

Extending this argument to arbitrary incentive-compatible outcomes (Propositions 1 and 2 and Corollary 4) is essentially equivalent to characterizing the conditions under which there exists an appropriate default decision. The potential difficulty is that either the default decision will not pose a sufficient threat for the agent or that for some proposals the principal will find the default decision more attractive and will therefore veto the decisions that should have been approved.³ Proposition 3 shows that the natural choice of the default decision is the most extreme decision that would be implemented on the equilibrium path: in this case the principal finds it optimal to approve every proposal expected on the equilibrium path and veto everything else.

In a more specialized environment, Lemma 1 characterizes the set of incentive-efficient outcomes. These outcomes can be implemented through veto-based delegation. Corollary 5 then shows that a smaller conflict of preferences between the agent and the principal results in a more extreme default decision for which veto-based delegation implements these efficient outcomes.

I also study which other mechanisms with imperfect commitment can implement efficient outcomes in this environment. Proposition 5 reveals that an efficient outcome can be implemented through a mechanism with imperfect commitment only if in this mechanism the principal has *almost no* ex-post discretion over the agent's decisions: the set of decisions that she can take after disapproving the agent's proposal should contain at most *one*, the most extreme, decision from the set of decisions that should be implemented on the equilibrium path. If the principal can take more than one such decision, she would excessively overrule the agent's proposals to correct for the conflict of preferences ex-post, which will negatively affect the agent's incentives and decrease ex-ante payoffs of both parties.

The analysis in this paper complements the results in Gilligan and Krehbiel [7], Krishna and Morgan [10], and Martin [12] that veto-based delegation (with a fixed default decision) may be superior to other forms of legislative rules and the result in Dessein [4] that full delegation is often superior to communication. This paper further extends this literature by showing that if the principal can select a default decision, veto-based delegation need not be inferior to any other decision making arrangement. This finding also clarifies the reasons for the conflicting conclusions of Dessein [4], who has shown that full delegation dominates veto-based delegation, and Marino [11], who presents a model in which the opposite is true.

Certainly, the principal who has absolute commitment power could implement a desired outcome through mechanisms other than veto-based delegation such as direct message mechanisms and constrained delegation. Section 8 discusses advantages and disadvantages of veto-based delegation over these alternative mechanisms.

The remainder of the paper is organized as follows: Section 2 presents an example. Section 3 introduces the model. Sections 4 and 5 derive the main results. Section

³In a similar model Melumad and Shibano [13] show that a subset of optimal incentive-compatible outcomes can be implemented under a form of veto-based delegation. Proposition 2 generalizes their result to other specifications of payoffs and distributions of private information, and to all incentive-compatible allocations rather than a subset of optimal allocations.

6 characterizes the efficient incentive-compatible outcomes, presents a comparative statics result about the default decision, and analyzes the optimal amount of ex-post discretion. In Section 7, I compare my results with the analysis of veto-based delegation and full delegation in Dessein [4]. Section 8 concludes.

2 An Example

There is a principal (she) who must take a decision $p \in \mathbb{R}$ and an agent (he) who has private information ω (his type) drawn from a uniform distribution on $[0, 1]$. The payoff of the principal is $u(p, \omega) = -(p - \omega)^2$, that of the agent is $u_a(p, \omega) = -(p - (\omega + b))^2$, where $b > 0$ is the bias of the agent. This is the setting of the main example in Crawford and Sobel [3] which has been used extensively in the literature on communication and delegation.

Holmström [8] points out that in this environment every feasible outcome can be viewed as an equilibrium of constrained delegation, where the agent is allowed to choose freely from a specified set of lotteries over decisions. This result is a twin of the Revelation Principle for the one-agent case. Consider an arbitrary game: In equilibrium, the actions of the agent result in a lottery over the final decisions. Instead of playing this game, one can ask the agent to pick a lottery directly by allowing him to choose from the set of lotteries present on the equilibrium path of the original game. Thus, without loss of generality we can concentrate on the games of constrained delegation.

Our goal is to see which outcomes of constrained delegation can be replicated through veto-based delegation. In this game, the principal selects the default decision, p_0 , the agent makes a proposal p , and, finally, the principal either approves or vetoes it. In the case of a veto, the default p_0 is implemented. The solution concept is the Perfect Bayesian equilibrium.

The difference between constrained delegation and veto-based delegation lies in the amount of commitment power given to the principal. Under constrained delegation the principal commits to approving some proposals and prohibiting others regardless of the information she infers from the agent's behavior. In contrast, under veto-based delegation, the principal should optimally make her approval decision given her *updated* beliefs.

Consider a case of constrained delegation whereby the agent is allowed to choose from $\mathcal{P} = [0, 1 - b]$. Because the agent's payoff function is a quadratic loss function in the difference between his most preferred alternative $p_a^I = \omega + b$ and the actual choice p , the agent will select the alternative from $\mathcal{P}$ that is the closest to p_a^I . In Figure 1 the bold line represents the outcome of the agent's decision problem. He chooses his most preferred alternative if $\omega \in [0, 1 - b]$ and $p = p_0 = 1 - b$ otherwise.

To replicate this outcome through veto-based delegation, set the default decision to $p_0 = 1 - b$. There follows an equilibrium in which the agent proposes $p = \omega + b$ for $\omega \in [0, 1 - b]$ and $1 - b$ otherwise, and the principal approves any proposal on the equilibrium path, $p \leq 1 - b$. She vetoes any proposal off the equilibrium path, $p > 1 - b$, which leads to the default decision $p_0 = 1 - b$.

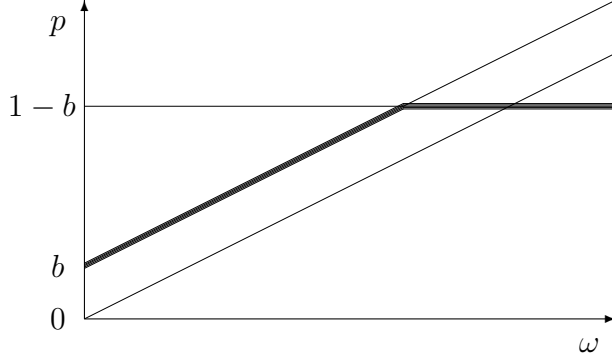


Figure 1: Constrained delegation with $\mathcal{P} = [0, 1 - b]$. The horizontal axis represents ω and the vertical axis represents p . The agent's bias is $b > 0$, ω is uniform on $[0, 1]$. The ideal decisions for the principal and the agent are correspondingly $p^I = \omega$ and $p_a^I = \omega + b$. The bold curve represents the equilibrium outcome: the agent chooses the alternative closest to his ideal.

On the equilibrium path, the principal infers ω when $p \in [b, 1 - b[$. She approves the proposal because the positive bias of the agent implies that the proposed decision p is better than the default decision p_0 . Off the equilibrium path, the principal's veto is optimal if, for example, she believes that $\omega = 1 - b$, in which case $p_0 = 1 - b$ is her ideal. The agent's behavior is optimal since he has effectively the same choices as in the original problem.

Thus, by setting p_0 equal to the highest decision taken on the equilibrium path under constrained delegation, the principal can replicate its outcome through veto-based delegation. In fact, for these specific preferences of the principal and the agent, a similar construction can be applied to any feasible outcome.

However, there are situations in which outcomes of constrained delegation cannot be achieved through veto-based delegation. For instance, assume that $u(p, \omega) = -(p - \omega)^2$ and $u_a(p, \omega) = -(p - (\omega/2 + 1/4))^2$ and allow the agent to choose any alternative from $\mathcal{P} = [1/3, 2/3]$, see Figure 2. (The principal will never actually select this mechanism; still, this example aids in illustrating why veto-based delegation may be unable to reproduce some outcomes.) In this case, the agent will always choose his most preferred alternative $p = \omega/2 + 1/4$ for $\omega \in [1/6, 5/6]$, $p = 1/6$ for $\omega < 1/6$ and $p = 2/3$ otherwise. It is impossible to find a default decision p_0 that allows replicating this outcome through veto-based delegation. Even if the agent follows the same strategy, when $p_0 \in [1/3, 2/3]$ the principal will find it optimal to intervene and veto some of the proposals on the equilibrium path. On the other hand, if $p_0 \in]-\infty, 1/3[\cap [2/3, +\infty[$, the agent will find it optimal to propose an alternative outside of $\mathcal{P}$ regardless of whether the principal will veto it. In either case, the original outcome cannot be achieved. Proposition 2 in Section 5 describes necessary and sufficient conditions under which any feasible outcome can be implemented through veto-based delegation.

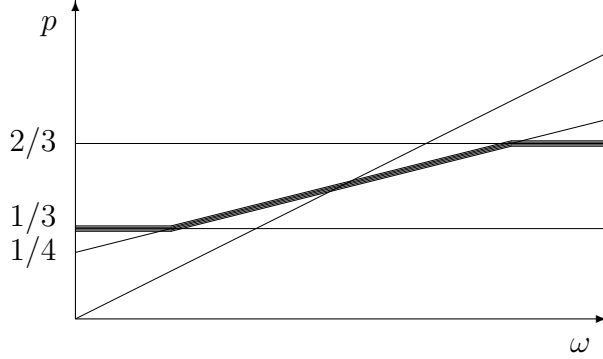


Figure 2: Constrained delegation with $\mathcal{P} = [1/3, 2/3]$. The horizontal axis represents ω and the vertical axis represents p . The agent's type ω is uniformly distributed on $[0, 1]$. The ideal decisions for the principal and the agent are correspondingly $p^I = \omega$ and $p_a^I = \omega/2 + 1/4$. The bold curve represents the equilibrium outcome: the agent chooses the alternative closest to his ideal.

3 The Model

There is a principal (she) and an agent (he). The agent has one-dimensional private information $\omega \in \mathbb{R}$, called a state of the world or a type of the agent. The principal's prior beliefs about ω are represented by a probability measure $\mu_\omega(\cdot)$ with a convex support Ω (notice that this allows for mass points, but not discrete support). The principal holds the rights over a convex set of possible decisions $\mathbb{P} \subseteq \mathbb{R}$. The payoffs of the principal and the agent are $u(p, \omega)$ and $u_a(p, \omega)$, where p is the decision taken. Both $u(\cdot, \omega)$ and $u_a(\cdot, \omega)$ are symmetric unimodal functions: they achieve their unique global maximums correspondingly at $p = p^I(\omega)$ and $p = p_a^I(\omega)$ and strictly decrease as p moves away from p^I and p_a^I .⁴ The functions $p^I(\cdot)$ and $p_a^I(\cdot)$, $p^I(\cdot) \neq p_a^I(\cdot)$ are continuous, bounded, non-decreasing, and satisfy $\mu(\Omega_\omega) = 0$ for $\omega = \{\omega' | p_a^I(\omega') = p_a^I(\omega)\}$, i.e., the probability of the states in which the most preferred decisions for the agent coincide is zero. Finally, $p^I(\cdot)$ is such that $\mathbb{P} = \{p | p^I(\omega) = p, \omega \in \Omega\}$.⁵

As discussed in Section 2 (see also Holmström [8] and Alonso and Matouschek [1]), due to the fact that in any game the agent is always free to choose any behavior regardless of his type, the knowledge of the final decisions taken on the equilibrium path is sufficient to characterize the equilibrium outcome; the specific content of the actions available to the agent is irrelevant. Therefore, a *mechanism* (with perfect

⁴A unimodal function can be viewed as a parametrization of a strictly quasiconcave monotonic utility function defined on a two dimensional space (e.g., quasilinear preferences) whereby choices are subject to a budget constraint. Unimodality is somewhat stronger than single-peaked preferences because it rules out indifference among decisions. Also, the assumption of symmetry of the payoffs is standard in the literature. It is convenient for our analysis but the results can be modified to account for asymmetric cases.

⁵This is also a standard (implicit) assumption in the literature. Even if it does not hold, one can redefine the set of feasible decisions as $\mathbb{P}' = \mathbb{P} \cap \{p | p = p^I(\omega)\}$. The only outcomes that may be potentially lost are strictly suboptimal from the principal's perspective.

commitment) is a set $\mathcal{P} \subseteq \mathbb{P}$.⁶ It induces a decision problem in which the agent receives his private information ω and chooses a decision $p \in \mathcal{P}$ to maximize his payoff. A mechanism is *ex-ante incentive-efficient* if it generates a Pareto efficient pair of the ex-ante payoffs (i.e., there is no mechanism that yields higher expected payoffs to the principal and the agent).⁷ It is assumed that the agent always participates in a mechanism.⁸

I also consider mechanisms in which the principal cannot ex-ante commit to a set of allowed decisions. Instead, after observing the agent's decision, the principal can intervene and change it. What the principal can commit to is the set of decisions available to her when overruling the agent's choice. Hence, a *mechanism with imperfect commitment* is a restricted set $\mathbb{P}_0 \subseteq \mathbb{P}$. It induces a game between the agent and the principal in which the agent proposes a decision $p \in \mathbb{P}$, and the principal either approves it or overrules it with a decision from $\mathbb{P}_0$.⁹ The solution concept for this game is the Perfect Bayesian equilibrium. That is, the agent's proposal and the principal's decision must be optimal given their beliefs, which, in turn, are required to be Bayesian whenever possible.

This definition of mechanism allows for communication, veto-based delegation, and full delegation, among other games. Under communication, the principal is free to make any decision after the agent's proposal, which is equivalent to $\mathbb{P}_0 = \mathbb{P}$. Under full delegation, the principal commits not to reverse the agent's decisions, which is equivalent to $\mathbb{P}_0 = \emptyset$. Finally, under veto-based delegation, if the agent's proposal is vetoed, then a fixed default decision p_0 is implemented. Here, the restricted decision set $\mathbb{P}_0 = \{p_0\}$ is a singleton.

An outcome is said to be *implemented through veto-based delegation* if there is some game of veto-based delegation in which this outcome is an equilibrium. The following two sections describe the conditions under which a given outcome can be implemented through veto-based delegation, but not necessarily as a unique equilibrium.

4 Imperfect Commitment

This section generalizes the observation that an outcome of communication can be implemented through veto-based delegation. Consider an arbitrary mechanism with *imperfect commitment*. Unless it is full delegation, in the equilibrium of this mechanism the principal chooses $p' \in \mathbb{P}_0$ following some (possibly, out-of-equilibrium) proposal by the agent. The outcome of this mechanism can be implemented through veto-based delegation. In order to do so, set $p_0 = p'$ and construct the equilibrium strategies by modifying the strategies in the original game as follows: If in the original

⁶Following the literature we consider only deterministic mechanisms. This may potentially contain some loss of generality since stochastic mechanisms may outperform deterministic.

⁷This concept of efficiency is defined in Holmström and Myerson [9].

⁸Alternatively, one can impose assumptions on the preferences such that efficient mechanisms generate a higher payoff for the agent than the decision the principal would take based solely on her prior beliefs.

⁹This is not the most general specification of a mechanism with imperfect commitment. See Bester and Strausz [2], who study the validity of the Revelation Principle under imperfect commitment.

game the agent's proposal was rejected, have the agent propose $p \in \mathbb{P}_0$, which was implemented instead of his proposal. Otherwise, his strategy is the same as before. The principal's strategy is to reject any proposal rejected in the original game and approve everything else. Finally, when an off-the-equilibrium-path proposal is vetoed, endow the principal with the belief that $\omega = \omega'$, where $p_0 = p^I(\omega')$; otherwise let her have the same beliefs as in the original game.

Certainly, if these strategies and beliefs are equilibrium then the outcomes are the same in both games. It is immediately clear that the principal's behavior off the equilibrium path is optimal given her beliefs. It is also easy to see the optimality of approving any proposal on the equilibrium path. First, after observing $p \notin \mathbb{P}_0$, the principal's beliefs are the same as in the original game and since it was optimal to approve the proposal before, it must be optimal to do so now. Second, for $p \in \mathbb{P}_0$, it must be that in the original game the agent's proposal was either p , in which case it was approved, or $p^* \neq p$, in which case it was vetoed with p . It follows that for all ω that lead to p , the principal prefers p to p_0 .

The agent's behavior is also optimal. Under veto-based delegation the agent achieves the same outcome as in the original game, while any deviation available to him now was also present in the original game. Therefore, the optimality of his original behavior implies that he does not have a profitable deviation.

We have

Proposition 1. *Any outcome of any mechanism with imperfect commitment, with the exception of full delegation, can be implemented through veto-based delegation.*

Proof. Let $p^A(\cdot)$ and $d^V(\cdot)$ be the equilibrium strategies under imperfect commitment and $P^A = \{p | p = d^V(p^A(\omega))\}$ be the set of implemented decisions. Choose any p' (it exists) such that $p' = d^V(p)$ for some p and set $p_0 = p'$. In the game of veto-based delegation define the strategies of the players by $p^a(\omega) = d^V(p^A(\omega))$ and

$$d^v(p) = \begin{cases} p, & p \in P^A; \\ p_0, & \text{otherwise.} \end{cases}$$

These strategies replicate the original outcome. For proposals off the equilibrium path, $p \notin P^A$, set the principal belief to be any ω such that $p_0 = p^I(\omega)$.

These strategies and beliefs are PBE. Consider the agent. Any deviation to $p \in P^A$, $p \neq p^a(\omega)$ was available in the original game and therefore cannot be profitable. Any deviation to $p \notin P^A$ results in p_0 ; it was also available before and hence is not profitable. For the principal, vetoing any $p \notin P^A$ is optimal due to the manner of construction of the out-of-equilibrium beliefs. It is optimal to approve proposals $p \in P^A$, since it was optimal to implement this decision rather than p_0 in the original game. $\square$

In particular,

Corollary 1. *Any outcome of communication can be implemented through veto-based delegation.*

Remark: full delegation. One can trivially implement full delegation through veto-based delegation by choosing the default decision such that the principal will never use it - if it is possible to find such a decision. Strictly speaking, under the assumption of our model that for any $p \in \mathbb{P}$ there exists $\omega \in \Omega$ such that $p = p^I(\omega)$, such implementation might be difficult. One can show that veto-based delegation can implement full delegation if and only if there exists ω_0 such that (1) the most preferred decisions of the agent and the principal coincide $p_a^I(\omega_0) = p^I(\omega_0)$, (2) if $p_a^I(\omega) \geq p^I(\omega_0)$ then $p^I(\omega) \geq (p^I(\omega_0) + p^I(\omega))/2$, and (3) if $p_a^I(\omega) \leq p^I(\omega_0)$ then $p^I(\omega) \leq (p^I(\omega_0) + p^I(\omega))/2$.

5 Perfect Commitment: The Veto-Power Principle

The previous section has shown that the principal can reproduce outcomes of mechanisms with imperfect commitment through veto-based delegation by asking the agent to propose the alternatives chosen in the original mechanism and promising to veto any deviations. The difficulty that arises if one attempts to adapt this result to the cases with perfect commitment is that the principal might be willing to veto some decisions that appear on the equilibrium path in the original mechanism. This is not an issue under imperfect commitment because in those games, by construction, every decision taken on the equilibrium path is optimal given the principal's beliefs. In contrast, if the principal has perfect commitment, she can, in fact, commit not to overrule some of the decisions. This section is devoted to characterizing the conditions under which any outcome induced by an arbitrary mechanism without monetary transfers can be implemented through veto-based delegation.

As shown in an example in Section 2, a sufficient condition that allows implementing a large range of incentive-compatible outcomes through veto-based delegation is that the agent is always (weakly) biased in the same direction, i.e., $p^I(\cdot)$ and $p_a^I(\cdot)$ do not intersect. At the same time, the example in Figure 2 points to a condition under which there are outcomes that cannot be implemented through veto-based delegation: the agent is biased towards higher decisions for low ideal decisions of the principal, but becomes more conservative and prefers lower decisions when the principal's ideal decisions are high. I call preferences for which this condition does not hold regular.

Definition 1. *The preferences are called regular if there do not exist $\omega^-, \omega^+ \in \Omega$, $\omega^- < \omega^+$ such that $p_a^I(\omega^-) - p^I(\omega^-) > 0 > p_a^I(\omega^+) - p^I(\omega^+)$.*

The following proposition generalizes examples in Section 2 by showing that any incentive-compatible outcome can be implemented through veto-based delegation if and only if the preferences are regular. The proof proceeds along the lines of the arguments in Section 2. The most tedious part is to demonstrate the optimality of approving the agent's proposals for the principal, when the agent is biased towards lower decisions for small ω and towards higher decisions for high ω . It is done by showing that if for low (resp. high) default decisions the principal is tempted to

overrule some of the proposals that should be accepted, she will find it optimal to approve all equilibrium proposals given some high (resp. low) default decision.

Proposition 2. (Veto-Power Principle) *Any outcome of any mechanism with perfect commitment can be implemented through veto-based delegation if and only if preferences are regular.*

Proof. Sufficiency. Fix an outcome described by $f : \Omega \rightarrow \mathbb{P}$ with a (compact) set of taken decisions $P^f = \{f(\omega) | \omega \in \Omega\}$. In a game of veto-based delegation with $p_0 \in \{\sup P^f, \inf P^f\}$, define the strategies of the players by $p^a(\omega) = f(\omega)$ and

$$d^v(p) = \begin{cases} p, & p \in P^f; \\ p_0, & \text{otherwise.} \end{cases}$$

Clearly, these strategies replicate the original outcome. For proposals off the equilibrium path, $p \notin P^f$, set the principal belief to be any ω such that $p_0 = p^I(\omega)$.

With a proper choice of p_0 these strategies and beliefs become PBE. Consider the agent. Any deviation to $p' \in P^f$, $p' \neq f(\omega)$ was available in the original mechanism and therefore cannot be profitable. Any deviation to $p' \notin P^f$ results in p_0 . Since p_0 is either supremum or infimum (or maximum or minimum) of P^f , this deviation also cannot be profitable. For the principal, vetoing any $p \notin P^f$ is optimal by the manner of construction of the out-of-equilibrium beliefs.

Showing the optimality of approving proposals on the equilibrium path is somewhat more involved. If for all ω , $p^I(\omega) - p_a^I(\omega) \leq 0$, set $p_0 = \sup P^f$. There are two possibilities. First, when $p^I(\omega) \leq p \leq p_0$, the optimality of approval is immediate. Second, if $p < p^I(\omega)$ then the optimality of the agent's strategy and the symmetry of his payoff implies $p_a^I(\omega) - p \leq p_0 - p_a^I(\omega)$. Therefore $p^I(\omega) - p \leq p_0 - p^I(\omega)$, which makes approval optimal. Similarly, the symmetric case is where for all ω , $p^I(\omega) - p_a^I(\omega) \geq 0$ and $p_0 = \inf P^f$.

Now consider the case in which $p^I(\omega) - p_a^I(\omega)$ varies its sign. First, consider the case in which P^f is convex. Denote by ω^* the state at which $p^I(\omega^*) - p_a^I(\omega^*) = 0$. The regularity of preferences implies that this state is unique, and that $p_a^I(\omega) \geq p^I(\omega)$ for $\omega \geq \omega^*$ and $p_a^I(\omega) \leq p^I(\omega)$ otherwise. Let ω^- satisfy $p_a^I(\omega^-) = \inf P^f$ and $p^- = p^I(\omega^-)$. Similarly, let ω^+ satisfy $p_a^I(\omega^+) = \sup P^f$ and $p^+ = p^I(\omega^+)$. Finally, denote $p^* = (\sup P^f + \inf P^f)/2$ and set

$$p_0 = \begin{cases} \sup P^f, & p^I(\omega^*) \leq p^*; \\ \inf P^f, & \text{otherwise.} \end{cases}$$

Assume $p^I(\omega^*) \leq p^*$. For all $\omega \geq \omega^*$, the equilibrium proposal satisfies $p \in [p^I(\omega), p_0]$ and therefore $p_0 - p^I(\omega) \geq p - p^I(\omega)$ which makes approval optimal. For $\omega \leq \omega^*$, approving the proposal is optimal if $p^I(\omega) - (p + p_0)/2 \leq 0$. The left hand side of this inequality achieves its supremum either at ω^- or ω^* . Thus, the proposal should be approved since $p^I(\omega^-) - (p_a^I(\omega^-) + p_0)/2 \leq p^I(\omega^*) - p^* \leq 0$.

Consider now the case in which P^f is not convex. That is, there are one or more intervals (p_1, p_2) , $\min P^f \leq p_1 < p_2 \leq \max P^f$, excluded from P^f . If $p^I(\omega^*) \leq p_1$ or $p_2 \leq p^I(\omega^*)$, nothing changes in the preceding argument. Assume now that $p_1 <$

$p^I(\omega^*) < p_2$. It is easy to see that $p_2 \in [p^I(\omega), p_0]$ and hence it should be approved. Let ω' be defined by $p_a^I(\omega') - p_1 = p_2 - p_a^I(\omega')$. Proposal p_1 will clearly be approved if for all ω such that $p^a(\omega) = p_1$, $p^I(\omega) \leq (p_2 + p_0)/2$. The left hand side of this inequality achieves its maximum at ω' . If $\omega' \geq \omega^*$, then $p^I(\omega') \leq p_a^I(\omega') = (p_1 + p_2)/2 \leq (p_0 + p_2)/2$. Otherwise, $\omega' \leq \omega^*$ and $p^I(\omega') \leq p^I(\omega^*) \leq p^* = (p_0 + \min P^f)/2 \leq (p_0 + p_2)/2$.

The proof for $p^I(\omega^*) \geq p^*$ is completely analogous.

Necessity. If preferences are not regular, then there exist ω^-, ω^+ , $\omega^- < \omega^+$ such that $p^I(\omega^-) < p_a^I(\omega^-)$ and $p^I(\omega^+) > p_a^I(\omega^+)$. Consider an outcome

$$f(\omega) = \begin{cases} p_a^I(\omega^-), & \omega \leq \omega^-; \\ p_a^I(\omega), & \omega^- < \omega < \omega^+; \\ p_a^I(\omega^+), & \omega^+ \leq \omega. \end{cases}$$

This outcome is incentive-compatible but cannot be implemented through veto-based delegation. In such a game, the agent would have to propose $p = p_a^I(\omega)$ for $\omega \in [\omega^-, \omega^+]$ which would allow the principal to infer ω . If $p_0 = \min P^f = p_a^I(\omega^-)$ or $p_0 = \max P^f = p_a^I(\omega^+)$, the principal would veto the proposals close to either $p_a^I(\omega^-)$ or $p_a^I(\omega^+)$. It is easy to see that no other decision could be used as a default. Either the default will not pose enough of a threat to the agent or the principal will find it profitable to veto some $p \in P^f$ to achieve her ideal decision. $\square$

This proposition does *not* imply that if the preferences are not regular, veto-based delegation cannot implement optimal outcomes, but rather that there will be some outcomes that could not be implemented. As evident in the example in Figure 2, such outcomes may be suboptimal. In fact, there, the optimal choice for the principal is to fully delegate decision-making to the agent. In that example, the outcome of full delegation can be implemented through veto-based delegation by setting $p_0 = 1/2$.

Certainly, if preferences are regular, then the optimal outcome can be implemented through veto-based delegation.

Corollary 2. *If preferences are regular, the outcome of the mechanism with perfect commitment that yields the highest payoff to the principal can be implemented through veto-based delegation.*

Another immediate consequence of this proposition follows from the fact that regularity of preferences and continuity of $p^I(\omega)$ and $p_a^I(\omega)$ imply that these functions intersect at most once.

Corollary 3. *If $p^I(\omega)$ and $p_a^I(\omega)$ intersect more than once, there are mechanisms whose outcomes cannot be implemented through veto-based delegation.*

In special cases the condition on preferences can be formulated in a somewhat simpler form. Often, the literature has studied quadratic payoffs with ideal decisions linear in ω . Therefore, I present

Corollary 4. *If $p^I(\omega) = a_p + b_p\omega$, $b_p > 0$, and $p_a^I(\omega) = a_a + b_a\omega$, $b_a > 0$, any outcome of any mechanism can be implemented through veto-based delegation if and only if either $p_a^I(\omega)$ and $p^I(\omega)$ never intersect or $b_p < b_a$.*

As can be seen in the example from Section 2, sometimes the decision that the principal has to set as a default is rather extreme. This is due to the fact that for moderate default decisions the principal will be tempted to overrule some proposals of the agent that need to be approved. Let $\underline{p}$ and $\bar{p}$ denote the lowest and the highest decisions that occur in some incentive-compatible outcome.

Proposition 3. *If an outcome can be implemented through veto-based delegation, then it can also be implemented through veto-based delegation with $p_0 \in \{\underline{p}, \bar{p}\}$. There exist outcomes which can only be implemented through veto-based delegation with $p_0 \in \{\underline{p}, \bar{p}\}$.*

Proof. The first part of the proposition has already been established in the proof of Proposition 2. The example given in Section 2 proves the second statement. In that example, the mechanism $\mathcal{P} = [0, 1 - b]$ can be implemented through veto-based delegation if and only if $p_0 = 1 - b$. If $p_0 > 1 - b$ the agent will deviate and offer $p = p_0$ for $\omega = p_0 - b$. If $p_0 < 1 - b$, the principal will veto the proposals around $p = p_0 + b$. $\square$

Although it is useful to know when *any* outcome can be implemented through veto-based delegation, a more pragmatic approach would be to characterize the conditions under which a *given* outcome can be implemented in this way. In general these conditions will depend on the precise functional form of prior beliefs and payoffs. However, one sufficient condition is that the implemented decision is always weakly bigger (or always weakly smaller) than the ideal decision of the principal.

Proposition 4. *An outcome $f(\omega_1)$ can be implemented through veto-based delegation if (1) for all ω $p^I(\omega) - f(\omega) \leq 0$ or (2) for all ω $p^I(\omega) - f(\omega) \geq 0$.*

Proof. Set $p_0 = \sup \mathcal{P}$, if for all ω $p^I(\omega) - f(\omega) \leq 0$, and $p_0 = \inf \mathcal{P}$ otherwise, and repeat the first part of the proof of sufficiency in Proposition 2. $\square$

The above results are obtained for deterministic mechanisms. In general, the Veto-Power Principle will not hold for mechanisms with perfect commitment that allow lotteries over decisions. After observing an agent's proposal, a risk-averse principal, for instance, will prefer to select a unique alternative rather than choose randomly among several of them.

Finally, the assumptions that $p^I(\cdot)$ and $p_a^I(\cdot)$ are bounded, continuous, and monotone or that Ω is convex are not crucial for the results in this and the previous sections. Even the assumption of the symmetry of the payoffs can be somewhat relaxed at the cost of complicating the regularity condition. However, without these assumptions it will be more difficult to characterize the necessary condition for the result in Proposition 2.

6 Optimal Amount of Ex-post Discretion

In this section I characterize the optimal amount of ex-post discretion $\mathbb{P}_0$ in the mechanisms with imperfect commitment. In order to do so, I first describe the set

of ex-ante incentive-efficient mechanisms and then show that their outcomes can be implemented under imperfect commitment only if $\mathbb{P}_0$ does not contain any decisions taken on the equilibrium path except the most extreme decision.

To make the analysis tractable I impose more structure on preferences and prior beliefs. It is assumed that ω is uniformly distributed on $[0, 1]$, $p^I(\omega) = \omega$, and $p_a^I(\omega) = b + \omega$, $b > 0$. In addition, $u(p, \omega) = u(p - \omega)$ and $u_a(p, \omega) = u_a(p - \omega - b)$ are strictly concave and symmetric around 0. I also restrict the set of possible decisions to be $\mathbb{P} = [b, 1]$.

Remark. Holmström [10] and Melumad and Shibano [15] have obtained the characterization of the optimal mechanism for the principal for a slightly more special case where the payoffs are quadratic and ω is distributed uniformly. Alonso and Matouschek [1] analyze the optimal delegation mechanism for continuous distributions of private information, single peaked payoff function of the agent, and a general quadratic payoff function of the principal. Martimort and Semenov [14] derive a condition on the prior beliefs for the optimal mechanism to be convex for the case of quadratic payoffs. Holmström mentions, without proof, that his result can be extended to more general payoff functions. It turns out that to do so one would have to make some additional assumptions similar to restricting the set of decisions to $\mathbb{P} = [b, 1]$.

Under our assumptions, the efficient mechanisms take a simple form, $\mathcal{P}_O = [b, p]$, $p \geq 1 - b$: the principal allows the agent to choose freely among moderate alternatives and prohibits all extreme decisions.¹⁰ In particular, the optimal mechanism for the agent has $p = 1$ and the optimal mechanism for the principal has $p = 1 - b$. The efficient mechanisms are convex because an exclusion of an interior set does not change the average decision taken in equilibrium but increases its variance, thus, hurting both parties. (This is exactly the reason why in Dessein [7] full delegation is often better than communication.)

Compare two mechanisms such that one of them has $[a_1, a_2] \in \mathcal{P}_1$ and the other has $a_1, a_2 \notin \mathcal{P}_2$, $a_1, a_2 \in \mathcal{P}_2$ (see, for example, Figures 1 and 3). In the first mechanism, the agent will choose $\omega + b$ for $\omega \in [a_1 - b, a_2 - b]$, while in the second he will choose a_1 for $\omega \in [a_1 - b, (a_1 + a_2)/2 - b]$ and a_2 for $\omega \in [(a_1 + a_2)/2 - b, a_2 - b]$. Due to the uniformity of distribution of ω the average realized decision and, therefore, the average distance between the ideal and the actual decision is the same in both mechanisms for either of the players. However, the variance of the realized decision is zero in the first mechanism and is strictly positive in the second mechanism. Therefore, the principal and the agent must be strictly worse off in the second mechanism due to strict concavity of their payoffs. This implies that the efficient mechanism should take the form $\mathcal{P}_O = [\underline{p}, \bar{p}]$.

Next, notice that the principal and the agent never want to exclude low decisions. If the lowest allowed decision $\underline{p}$ exceeds b then for $\omega \in [0, \underline{p} - b]$ the agent must choose $\underline{p}$, whereas both the agent and the principal would have been better off with a decision

¹⁰Without the restriction that $\mathbb{P} = [b, 1]$, a principal who is not too risk-averse could benefit from excluding a set of decisions around b since it would decrease the average realized distance between the taken decision and her most preferred decision at the cost of additional variance. In this case, some efficient mechanisms could be of the form $\mathcal{P}_O = \{p_1\} \cup [\underline{p}, \bar{p}]$.

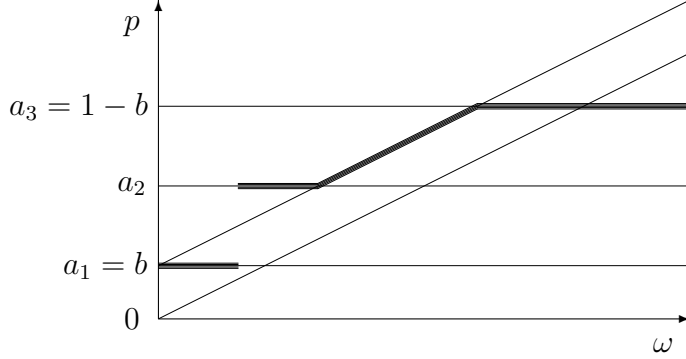


Figure 3: Constrained delegation with $\mathcal{P} = \{a_1\} \cup [a_2, a_3]$. The horizontal axis represents ω and the vertical axis represents p . The agent's bias is $b > 0$, ω is uniform on $[0, 1]$. The ideal decisions for the principal and the agent are correspondingly $p^I = \omega$ and $p_a^I = \omega + b$. The bold curve represents the equilibrium outcome: the agent chooses the alternative closest to his ideal.

of $\omega + b$. Hence, the optimal mechanism must have $\mathcal{P}_O = [b, \bar{p}]$.

Finally, let the mechanism be $\mathcal{P} = [b, \bar{p}]$. In this case, the agent will choose $p = \omega + b$ for $\omega \in [0, \bar{p} - b]$ and p otherwise. The principal obtains the decision which is b away from her most preferred decision for $\omega \in [0, \bar{p} - b]$ and the decision that is $|\bar{p} - \omega|$ away from her most preferred decision otherwise. Because of uniformity of the distribution of ω , the principal's payoff is maximized when $\bar{p} = 1 - b$ (see the proof of Lemma 1 for details). At the same time, the agent's payoff increases in p since higher values enable him to obtain decisions closer to his ideal. This allows us to conclude that every efficient mechanism must have $p \geq 1 - b$. Therefore,

Lemma 1. *The set of efficient mechanisms is $\mathbb{P}^E = \{[b, \bar{p}] | \bar{p} \geq 1 - b\}$. The highest payoff for the principal is achieved by $\mathcal{P}^P = [b, 1 - b]$ and the highest payoff for the agent is achieved by $\mathcal{P}^A = [b, 1]$.*

Proof. Let $\mathcal{P}_1$ and $\mathcal{P}_2$ be mechanisms that differ only in that $]a_1, a_2[\in \mathcal{P}_1$ and $]a_1, a_2[\notin \mathcal{P}_2$, where $a_1, a_2 \in \mathcal{P}_1, \mathcal{P}_2$. Then, the parties can achieve different payoffs in these two mechanisms only for $\omega \in \Omega^* =]a_1 - b, a_2 - b[$. For all of these states, in equilibrium of $\mathcal{P}_1$, the agent will select $p = \omega + b$. In equilibrium of $\mathcal{P}_2$, however, the agent will select $p = a_1$ for $\omega \in \Omega_-^* =]a_1 - b, (a_1 + a_2)/2 - b[$ and $p = a_2$ for $\omega \in \Omega_+^* = [(a_1 + a_2)/2 - b, a_2 - b[$. Therefore, conditional on Ω^* in these mechanisms the expected payoffs of the principal are $U(\mathcal{P}_1|\Omega^*) = u(b)$ and $U(\mathcal{P}_2|\Omega^*) = (U(\mathcal{P}_2|\Omega_-^*) + U(\mathcal{P}_2|\Omega_+^*))/2 = (\int_{\Omega_-^*} u(a_1 - \omega) d\omega + \int_{\Omega_+^*} u(a_2 - \omega) d\omega)/2 < u(b)$. Similarly, the expected payoffs of the agent are $U_a(\mathcal{P}_1|\Omega^*) = u(0)$ and $U_a(\mathcal{P}_2|\Omega^*) = (U_a(\mathcal{P}_2|\Omega_-^*) + U_a(\mathcal{P}_2|\Omega_+^*))/2 = (\int_{\Omega_-^*} u_a(a_1 - \omega - b) d\omega + \int_{\Omega_+^*} u_a(a_2 - \omega - b) d\omega)/2 < u(b)$. This implies that any ex-ante efficient mechanism must be convex.

Next, let us show that the lowest decision in any efficient mechanism is $\underline{p}_O = b$. The alternative would be to have $\underline{p}_O > b$. For $\omega \in [0, \underline{p} - b]$, the agent would choose

$p = \omega$ in the former case and $p = \underline{p}_O$ in the latter. Since $p = \omega$ is closer to the most preferred decisions of both parties, and for all other ω the chosen decision is the same in both mechanisms, it must be that $\underline{p}_O = b$.

It is left to prove that the highest decision in any efficient mechanism is $\bar{p}_O \geq 1 - b$. Since the efficient mechanism is convex and $\underline{p}_O = b$, in any such mechanism the decisions taken by the agent are

$$d(m(\omega)) = \begin{cases} \omega + b, & \omega < \bar{p}_O - b; \\ \bar{p}_O, & \text{otherwise.} \end{cases}$$

For a given $\bar{p}_O$ the expected payoff for the agent is $U_a(\bar{p}_O) = u_a(0)(\bar{p}_O - b) + \int_0^{1-\bar{p}_O+b} u_a(\omega) d\omega$, which increases in p . The expected payoff for the principal is

$$U(\bar{p}_O) = \begin{cases} u(b)(\bar{p}_O - b) + 2 \int_{1-b}^1 u(\omega) d\omega + \int_0^{1-\bar{p}_O-b} u(b + \omega) d\omega, & \bar{p}_O \leq 1 - b; \\ u(b)(\bar{p}_O - b) + \int_0^b u(\omega) d\omega + \int_0^{1-\bar{p}_O} u(\omega) d\omega, & \text{otherwise.} \end{cases}$$

It achieves its maximum at $\bar{p}_O = 1 - b$ and decreases afterwards. Therefore, the highest payoff for the principal is achieved by $\mathcal{P}^P = [b, 1 - b]$, the highest payoff for the agent by $\mathcal{P}^A = [b, 1]$, and the set of efficient mechanisms is $\mathbb{P}^E$. $\square$

Proposition 2 implies that outcomes of these mechanisms can be implemented through veto-based delegation. I now show that these outcomes can be implemented through other mechanisms with imperfect commitment if and only if, with the exception of one decision, the principal abandons the right to take decisions which appear on the equilibrium path and instead relies completely on the proposals of the agent: the set of decisions taken by the agent on the equilibrium path $\mathcal{P} = [b, \bar{p}]$ and the set of decisions available for the principal $\mathbb{P}_0$ must have one and only one decision in common, the largest decision taken on the equilibrium path, $\bar{p}$.¹¹

To see this, consider an efficient mechanism with $\bar{p} < 1$. Imagine that there is $p' \in \mathbb{P}_0 \cap \mathcal{P}$, $p' < \bar{p}$. Then the principal will use it to overrule proposals in $(p' - b, \min\{\bar{p}, p' + b\})$ to correct for the agent's bias. On the other hand, if $\mathbb{P}_0 \cap \mathcal{P} = \emptyset$, then there will be decisions taken that exceed $\bar{p}$.

Proposition 5. *An outcome of a mechanism $\mathcal{P} = [b, \bar{p}]$, $1 > \bar{p} \geq 1 - b$ can be implemented through a mechanism with imperfect commitment if and only if $\mathbb{P}_0 \cap \mathcal{P} = \bar{p}$. The mechanism $\mathcal{P}^A = [b, 1]$ can be implemented through a mechanism with imperfect commitment if and only if $\mathbb{P}_0 = 1$ or $\mathbb{P}_0 = \emptyset$.*

Proof. Clearly, the outcome of $\mathcal{P} = [b, 1]$ could be implemented by full delegation with $\mathbb{P}_0 = \emptyset$.

Next, in any equilibrium of a mechanism with $\mathcal{P} = [b, \bar{p}]$ the agent chooses

$$p = \begin{cases} \omega + b, & \omega < \bar{p} - b; \\ \bar{p}, & \text{otherwise.} \end{cases}$$

¹¹With the trivial exception of the outcome of full delegation, which can be implemented both with $\mathbb{P}_0 = \bar{p} = 1$ and $\mathbb{P}_0 = \emptyset$.

and for $p < \bar{p}$ the principal's expected payoff from decision p' after a proposal p is

$$\mathbb{E}u(p'|p) = u(p', p - b).$$

If there exists $p_0 \in \mathbb{P}_0 \cap \mathcal{P}$, $p_0 \neq \bar{p}$, then $p_0 < \bar{p}$. In this case, for any proposal $p \in]p_0, p_0 + b]$,

$$\mathbb{E}u(p|p) = u(p, p - b) < u(p_0, p - b) = \mathbb{E}u(p_0|p).$$

Therefore, the principal is better off overruling all such p with p_0 , which would destroy the desired outcome.

Now consider the case in which $\mathbb{P}_0 \cap \mathcal{P} = \emptyset$, $\mathbb{P}_0 \neq \emptyset$ and $\bar{p} < 1$. The only possibility to support the desired outcome in equilibrium is for the principal to overrule all proposals $p > \bar{p}$ with some $p_0 \in \mathbb{P}_0$. However, because in this case $\bar{p} < p_0 \leq 1$, there exists ϵ , $0 < \epsilon < p_0 - \bar{p}$ such that for $\omega \in]p_0 - b - \epsilon, p_0 - b]$,

$$u_a(\bar{p} - \omega - b) < u_a(p_0 - \omega - b)$$

Hence, the agent is better off proposing some $p > \bar{p}$ that will be vetoed by p_0 rather than proposing $\bar{p}$ for all such ω , which would destroy the desired outcome. $\square$

Lemma 1 and Proposition 5 yield a comparative statics result:

Corollary 5. *The default decision for which veto-based delegation implements the optimal outcome for the principal, $p_0 = 1 - b$, becomes more extreme as the bias of the agent b becomes smaller.*

For the setting with more general distribution of private information, non-constant bias, and quadratic payoff of the principal, Propositions 3, 4, 5, and 6 in Alonso and Matouschek [1] establish the conditions under which, similarly to the result in Lemma 1, the optimal mechanisms for the principal take the form of a single interval. Proposition 5 and Corollary 5 would extend to their environment under the condition that the bias of the agent has a constant sign.

7 Veto-based Delegation vs. Full Delegation

Proposition 7 in Dessein [4] compares the relative performance of full delegation and veto-based delegation for fixed default decisions p_0 in settings where the agent has a constant positive bias, ω is distributed uniformly and payoff functions are concave. Let $b(p_0) = (1 - p_0)/3$. His result can be restated as:

- (i) If $b \leq b(p_0)$, full delegation strictly dominates veto-based delegation;
- (ii) If $b > b(p_0)$, veto-based delegation dominates full delegation if the principal is not too risk-averse.

This result is related to Proposition 5, which states that veto-based delegation is the only mechanism with imperfect commitment that achieves the incentive-efficient outcomes $\{[b, \bar{p}] | 1 > \bar{p} \geq 1 - b\}$. In particular, this means that full delegation cannot achieve the principal's favorite outcome for any degree of risk-aversion, and is therefore payoff-inferior to veto-based delegation.

More generally, one can show that for all $p_0 > 1 - b$ the principal is strictly better off under veto-based delegation than under full delegation: it pays to exclude high decisions for any degree of risk-aversion of the principal.

Example. If $p_0 \geq 1 - b$, ω is uniformly distributed on $[0, 1]$, $p^I(\omega) = \omega$, $p_a^I(\omega) = b + \omega$, $b > 0$, $u(p, \omega) = u(p - \omega)$ and $u_a(p, \omega) = u_a(p - \omega - b)$ are strictly concave and symmetric around 0, and $\mathcal{P} = [b, 1]$, there is an equilibrium under veto-based delegation that yields to the principal a strictly higher payoff than does full delegation.

Proof. Consider the following equilibrium under veto-based delegation (Part (ii) of Lemma 4 in Dessein [4]). The agent proposes his ideal decision for $\omega \leq p_0 - b$ and p_0 otherwise. The principal approves everything below p_0 and vetoes everything above it. The principal's beliefs are Bayesian for the proposals on the equilibrium path and $\omega = p_0$ otherwise (notice that these beliefs satisfy the Strict Equilibrium Dominance for $p_0 \geq 1 - b$). The agent's behavior is optimal given that the principal vetoes anything above p_0 . The principal's veto is optimal given her beliefs. Due to the fact that the agent's bias is positive, the principal cannot benefit from the veto of any proposal off the equilibrium path. When $p_0 > 1 - b$, this equilibrium yields $u(b)(p_0 - b) + \int_0^b u(\omega) d\omega + \int 1 - p_0 u(\omega) d\omega$ to the principal. On the other hand, under full delegation the payoff is $u(b)$. Clearly, veto-based delegation yields a higher payoff than does full delegation.

8 Discussion: veto-based delegation vs. other mechanisms

Veto-based delegation is a commonly observed decision-making arrangement. In organizations, the decisions of employees often require approval of higher-level management. In many legal institutions one of the decision-making parties has veto rights. In this paper I argue that veto-based delegation is an attractive decision mechanism because with a proper choice of the default decision it achieves the optimal separation of decision initiation and decision control (Fama and Jensen [5]): the principal encourages the agent to use his knowledge by delegating him *formal* rights to initiate and implement decisions, while she prevents the agent from behaving opportunistically by keeping the right to block his decision.

In addition to veto-based delegation, incentive-compatible outcomes can be implemented by at least two other types of mechanisms: message mechanisms and constrained delegation. In a message mechanism, the principal specifies a set of messages and a function that maps the message sent by the agent into a decision. Under constrained delegation, the principal specifies the set of decisions from which the

agent chooses a decision. Veto-based delegation can do as well as any of these mechanisms, but clearly it cannot achieve strictly more. Moreover, often (but not always) these mechanisms implement the desired outcome as a unique equilibrium, whereas veto-based delegation typically has several equilibria.¹²

One of the advantages of veto-based delegation over these alternative mechanisms is its simplicity: veto-based delegation requires specifying only a default decision, whereas the other mechanisms must specify *all* feasible decisions. The fact that determining and describing the entire set of decisions can be prohibitively costly has long been recognized: Justice White noted in his dissenting opinion in *INS v. Ghadha*, 462 U.S. 919, (1983) that without the legislative veto Congress would be:

faced with a Hobson’s choice: either to refrain from delegating the necessary authority, leaving itself with a hopeless task of writing laws with the requisite specificity to cover endless special circumstances across the entire policy landscape, or in the alternative, to abdicate its lawmaking function to the Executive Branch and independent agencies.¹³

There could also be other reasons that favor veto-based delegation. For instance, in case of constrained delegation the principal authorizes the agent to take decisions on her behalf. In these situations the doctrine of apparent authority implies that the principal might be bound by the actions of the agent even if they were not authorized.¹⁴ In *the American Society of Mechanical Engineers (ASME) v. Hydrolevel Corp*, 456 U.S., 556, (1982), a predisposed member of ASME issued an *unofficial* opinion declaring a Hydrolevel’s product unsafe. The Hydrolevel’s competitor used this opinion to discourage sales of the Hydrolevel’s product, which eventually put Hydrolevel out of business. The Supreme Court found ASME liable on the principle of apparent authority. After this incident, ASME had mandated that all its opinions should be reviewed and approved by someone not involved in their preparation.

As mentioned above, one difficulty with veto-based delegation is that the freedom in constructing beliefs off the equilibrium path creates multiplicity of equilibria. However, there is some evidence that in reality the principal and the agent may be able to coordinate to play a specific equilibrium. For instance, Fisher [6], p. 289, presents an example in which NASA and the US Congress document their expectations about how the law appropriating funds for NASA should be carried out.

Finally, the results in this paper can also be viewed from a different perspective. In many situations veto-based delegation is not an outcome of a careful mechanism design but rather a myopic attempt of a principal to deal with an ever increasing

¹²There could be multiple equilibria in message mechanisms and in constrained delegation if in the equilibrium implementing the intended outcome the agent is indifferent between several messages or decisions.

¹³Quoted from Fisher [6]. The US Supreme Court concluded in *INS v. Ghadha* that legislative veto is unconstitutional. The expedited congressional oversight procedures were later developed as an alternative form of veto-based delegation not subject to this ruling.

¹⁴The principle of apparent authority states that the principal is liable for the acts of her agent even if the agent does not have an authority to perform this action but appears to the third party as if he were granted such authority by the principal.

number of tasks: delegate to the agent and retain the right to block his decisions. This paper suggests that when such delegation takes place its success will hinge on the proper choice of the default decision.

References

- [1] ALONSO, R., AND MATOUSCHEK, N. Optimal delegation. Northwestern University, 2005.
- [2] BESTER, H., AND STRAUZ, R. Contracting with imperfect commitment and the revelation principle: The single agent case. *Econometrica* 69 (2001).
- [3] CRAWFORD, V. P., AND SOBEL, J. Strategic information transmission. *Econometrica* 50 (1982).
- [4] DESSEIN, W. Authority and communication in organizations. *Review of Economic Studies* 69 (2002).
- [5] FAMA, E. F., AND JENSEN, M. C. Separation of ownership and control. *Journal of Law and Economics* 26 (1983).
- [6] FISHER, L. The legislative veto: Invalidated, it survives. *Law and Contemporary Problems* 56 (1993).
- [7] GILLIGAN, T. W., AND KREHBIEL, K. Collective decision-making and standing committees: An informational rationale for restrictive amendment procedures. *Journal of Law, Economics, and Organization* 3 (1987).
- [8] HOLMSTRÖM, B. On the theory of delegation. In *Bayesian Models in Economic Theory*, M. Boyer and R. E. Kihlstrom, Eds. North-Holland, 1984, pp. 115–41.
- [9] HOLMSTRÖM, B., AND MYERSON, R. B. Efficient and durable decision rules with incomplete information. *Econometrica* 51, 6 (1983), 1799–1819.
- [10] KRISHNA, V., AND MORGAN, J. Asymmetric information and legislative rules: Some amendments. *American Political Science Review* 95 (2001).
- [11] MARINO, A. M. Delegation versus veto in organizational games of strategic communication. Marshall School of Business, 2005.
- [12] MARTIN, E. M. An informational theory of the legislative veto. *Journal of Law and Economics* 13, 2 (1997), 319–43.
- [13] MELUMAD, N. D., AND SHIBANO, T. The securities and exchange commission and the financial accounting standards board: Regulation through veto-based delegation. *Journal of Accounting Research* 32 (1994).