

Regulating Information*

PEDRO GONZALEZ-FERNANDEZ[†] STEFAN TERSTIEGE[‡] ELIAS TSAKAS[§]

Maastricht University

June, 2025

Preliminary draft

Abstract

Economic agents often seek to acquire information before taking an important action, but in many domains, gathering this information requires approval from a regulator. This paper develops a model in which an agent designs an experiment to inform his decision, but can only implement it if a regulator authorizes it ex ante. We characterize the agent's optimal experiment under this approval constraint and show that, whenever the regulator rejects full revelation, the agent strategically reduces informativeness in the states where their disagreement is least sensitive. We then extend the model to settings with multiple regulators, comparing sequential and collective approval mechanisms. The analysis yields predictions for how institutional structure shapes access to information, with applications to clinical trials, data privacy, and ethics boards.

Keywords Information Design · Regulation · Institutional Mechanisms

JEL Classification D83 · L51

*We are thankful to Andrew Ellis, Alexis Ghersegorin, Inés Moreno de Barreda, and Ludvig Sinander for helpful comments and suggestions. We also thank Ferdinand Pieroth and Eline Wijns for helpful discussions during the development of this paper.

[†]Homepage: <https://pgonzalezfernandez.com>; E-mail: p.gonzalezfernandez@maastrichtuniversity.nl

[‡]Homepage: [stefanterstiege/home](https://stefanterstiege.com); E-mail: s.terstiege@maastrichtuniversity.nl

[§]Homepage: <http://www.elias-tsakas.com/home.html>; E-mail: e.tsakas@maastrichtuniversity.nl

1 Introduction

Should an employer be allowed to test whether a candidate has a genetic condition? Should a police department target searches based on ethnicity-related data? Should a company be allowed to access users' private data to tailor prices or offers? These questions exemplify a growing tension across many domains: agents often want to acquire precise information before taking an action, but regulators may wish to restrict the use or collection of such information due to ethical, legal, or societal concerns.

In such settings, regulators are not only concerned with the actions agents take, but with the experiments they run to inform those actions. Importantly, these regulators often act as a representative of broader societal interests. For example, society may object to employers accessing genetic records, or to insurers pricing based on personal medical history. Likewise, individuals may oppose firms harvesting private data to tailor prices or services. These concerns reflect the broader insight that not all information is socially beneficial (Hirshleifer, 1971): while it may improve the agent's private decision-making, it can also undermine fairness, risk-sharing, or consumers' privacy (Hackmann et al., 2015; Agan and Starr, 2018; Galperti and Perego, 2023), and in turn, reduce overall social welfare. From this perspective, the regulator's role is to prevent precisely those experiments that might generate private benefits but impose broader social costs.

This type of institutional oversight is widespread—in domains ranging from clinical trials and criminal investigations to data privacy and employment screening¹—but has received limited attention in standard economic models. Existing frameworks in information design typically focus on settings where the experimenter acts alone or is constrained only by internal objectives. While recent work has begun incorporating external limitations, such as discriminatory concerns or privacy constraints (Onuchic and Ray, 2023; Strack and Yang, 2024), the role of formal institutional approval, where an external regulator is responsible to authorize the agent's information acquisition, remains understudied. Un-

¹For instance, Institutional Review Boards (IRBs) and FDA panels must approve clinical trials in advance; under the GDPR, data collection protocols may require prior consultation with supervisory authorities; and U.S. law limits pre-employment screening through regulations such as the Genetic Information Nondiscrimination Act. See U.S. Department of Health and Human Services (2022), 45 CFR §46.108; European Union (2016), GDPR Articles 64–65; and U.S. Equal Employment Opportunity Commission (2010).

derstanding how such regulatory constraints shape experimentation is both conceptually and practically important. It has direct implications for how firms design screening procedures, how companies access private data, or how researchers run experiments involving human subjects. This paper develops a simple and flexible framework to analyze this question.

We study a model in which an agent chooses between two actions whose payoffs depend on an unknown state of the world. Before acting, the agent can design an experiment to learn about the state, but it must be approved by a regulator who shares the same prior but values actions differently. The regulator evaluates the experiment *ex ante* and approves it only if it yields a higher expected payoff than the agent’s baseline action under the prior. This creates a constraint on the space of admissible experiments. The central tension is simple: the agent prefers the most informative experiment (i.e. the one that perfectly reveals the state) but the regulator may block it if it induces actions she finds undesirable in some states. In those cases, we find that the agent must design an experiment that leaves the regulator exactly indifferent between approval and rejection. As a result, when the agent and the regulator disagree, the regulator reaches exactly her objective: She is not made worse off by the experimenter; but not better off either. We then characterize the structure of optimal experiments under this approval constraint. The key insight is that the agent distorts the experiment selectively, reducing informativeness in states where his disagreement with the regulator is most salient. This yields a full characterization of optimal experiments along the Pareto frontier of feasible payoffs.

Example 1. Suppose an employer is deciding whether to hire a promising candidate for a long-term position. Based on the interview, the employer suspects there is a 20% chance the candidate may have cancer. He is risk-neutral and would prefer not to hire someone with a serious medical condition. Human Resources (HR), however, serves as a regulator in this context: it shares the same information as the employer but is concerned only with a fair evaluation based on the candidate’s qualifications, not her health status. Table 1 shows the payoffs to the employer (agent) and HR (regulator) under each action and state.

With no additional information, the employer chooses to hire. However, he would prefer to run a background check to learn the candidate’s health status before deciding. If the background check perfectly reveals whether the candidate has cancer, the employer will

		States	
		<i>Cancer</i>	<i>No Cancer</i>
Actions	<i>Hire</i>	$(-1, 1)$	$(1, 1)$
	<i>Not Hire</i>	$(0, 0)$	$(0, 0)$

Table 1: Payoffs for agent and regulator.

always decline to hire her when the condition is confirmed—resulting in lower expected utility for HR than the baseline of hiring without further information. Anticipating this, HR blocks the use of fully revealing tests. In response, the employer designs a background check that is deliberately noisy: fuzzy enough that he sometimes hires even if the candidate has cancer. By reducing informativeness in just the right way, he leaves HR exactly indifferent between approval and rejection. The experiment is thus distorted not uniformly, but selectively: informativeness is sacrificed precisely where HR’s and the employer’s incentives diverge.

While our baseline model involves a single regulator, many institutional settings involve more than a single regulator. In some cases, an agent can appeal a rejected proposal to higher authorities (e.g. IRBs and FDA in clinical trials); in others, a panel or board jointly decides whether to authorize an experiment (e.g. GDPR boards). To capture such structures, we extend the model to environments where multiple regulators must approve the agent’s information acquisition. Even when regulators share aligned preferences, the order and structure of the approval process play a critical role in shaping what kinds of experiments are ultimately allowed.

To analyze such settings, we introduce a comparative framework that formalizes how permissive or restrictive a regulator is (what we call *strictness*). This notion allows us to rank regulators by how tightly they constrain the agent’s ability to acquire information. When approval is sequential, such as in appeals, what matters most is who reviews the experiment last. If the final authority is highly restrictive, the agent must dilute the experiment enough to satisfy her, even if earlier reviewers were more permissive. When approval is collective, such as through committee voting, the pivotal regulator is the one in the middle. In this case, the implemented experiment reflects the preferences of the median reviewer, not the strictest or most lenient. We illustrate these insights through

applications to clinical trials, criminal investigations, data privacy, and medical review boards.

Related Literature. Concerns about harmful or sensitive information are well documented across a range of applied settings, but typically addressed in isolation. [McClellan \(2022\)](#) studies how ethics boards regulate clinical trials to limit health risks. [Strack and Yang \(2024\)](#) analyze how data protection rules constrain firms’ ability to extract private information under privacy-preserving detection constraints. [Agan and Starr \(2018\)](#) show that employment regulations, such as Ban-the-Box policies, restrict access to criminal records in hiring to reduce discrimination. These papers reflect a growing recognition that society may not want certain types of information to be acquired or used, even when doing so would improve decision-making ([Hirshleifer, 1971](#); [Morris and Shin, 2002](#); [Tirole, 2016](#)). What remains missing, however, is a unified framework to study how these concerns operate across domains, and how they shape the design of experiments and information acquisition. We address this gap by modeling the regulator as an institutional proxy for society. That is, one that evaluates experiments *ex ante* and blocks those that would reduce expected social welfare.

In canonical models of information design, such as [Kamenica and Gentzkow \(2011\)](#), a sender strategically designs experiments to influence a receiver’s action, typically under full autonomy. Subsequent work has explored settings in which the sender faces limitations—whether technological, informational, or ethical. [Doval and Skreta \(2024\)](#) study arbitrary restrictions on the set of feasible experiments, while [Ichihashi \(2019\)](#) examines how limits on the sender’s own information affect optimal persuasion. Privacy-motivated constraints also appear in [Ichihashi \(2020\)](#), where consumers restrict the flow of data to digital platform. While these papers consider rich constraint structures, they generally treat the limits as exogenous and internal to the sender’s environment. In contrast, we model regulation as an institutional approval constraint: the regulator evaluates the experiment *ex ante* and can veto its implementation. This external approval requirement alters the agent’s optimization problem, forcing him to design experiments that are deliberately less informative in ways that align with the regulator’s preferences.

Our paper also contributes to a growing literature on delegated experimentation and regulatory approval. Several studies model institutional review as an approval constraint or persuasion device: [Henry and Ottaviani \(2019\)](#) analyze sequential drug trials under

regulatory oversight, while [Hu and Sobel \(2022\)](#) interpret approval as a form of costly persuasion. Closest in spirit is [Wittbrodt and Yoder \(2025\)](#), who study how a principal can screen an agent’s private information by offering a menu of acceptable experiments. In contrast, we consider a symmetric-information environment where the agent designs a single experiment subject to external approval, and where approval operates as a binary constraint rather than a mechanism design menu. Methodologically, our paper is broadly related to the literature on sequential information design. [Doval and Ely \(2020\)](#), which examines how the timing of information release and commitment constraints affect persuasion outcomes. Finally, our extension to multiple regulators introduces comparative predictions across sequential and collective review institutions, providing insights into appeals processes and board-based approvals, and relating to institutional models of persuasion in political economy (e.g., [Alonso and Câmara, 2016](#); [Bardhi and Guo, 2018](#)).

The remainder of the paper is structured as follows: Section 2 introduces the model; Section 3 characterizes the optimal experiment under approval constraints, Section 4 extends the framework to analyze multiple regulators, and derives comparative predictions; Section 5 applies the results to institutional settings and compares sequential and collective approval processes; Section 6 concludes. All proofs are relegated to the appendix.

2 The Model

We consider a strategic interaction between an *agent*, who seeks to make an informed decision, and a *regulator*, who can approve or block the agent’s experiment. The agent chooses between two actions, $A = \{a_1, a_2\}$, whose consequences depend on the realized state $\omega \in \Omega$. The agent and the regulator share a common prior over states given by $\bar{\mu} \in \Delta(\Omega)$. The agent’s payoff from action a in state ω is $u(a, \omega)$, and the regulator’s payoff is $v(a, \omega)$. We use the index $k \in \{1, 2\}$ to denote any available action a_k , while $a_g = \arg \max_a u(a, \omega)$ and $a_r = \arg \max_a v(a, \omega)$ for $g, r \in \{1, 2\}$ denote the agent’s and the regulator’s most preferred actions in state ω , respectively. That is, a_g and a_r serve as shorthand for the state-dependent optimal choices $a_g(\omega)$ and $a_r(\omega)$; we omit the argument when the state is clear from context.

The agent can run an *experiment* to obtain information before choosing an action. An experiment is a mapping $\sigma_k : \Omega \rightarrow [0, 1]$ for $k = 1, 2$, where $\sigma_k(\omega)$ is the probability of receiving signal a_k in state ω . Since the action space is binary, two signal realizations

suffice to describe any experiment without loss of generality. Accordingly, we set $\sigma_1(\omega) = 1 - \sigma_2(\omega)$. Thus, the posterior probability of state ω given that the agent receives signal a_k is determined by Bayes' rule:

$$\mu(\omega|a_k) = \frac{\bar{\mu}(\omega)\sigma_k(\omega)}{\sum_{\omega' \in \Omega} \bar{\mu}(\omega')\sigma_k(\omega')}. \quad (1)$$

We use σ as shorthand for an experiment and denote the set of all experiments by Σ . Without loss of generality, the agent follows the experiment's recommendation only if it is *incentive compatible*, meaning that after receiving signal a_k , the agent prefers action a_k . This imposes the following conditions:

$$\sum_{\omega} u(a_1, \omega) \mu(\omega|a_1) \geq \sum_{\omega} u(a_2, \omega) \mu(\omega|a_1), \quad \sum_{\omega} u(a_2, \omega) \mu(\omega|a_2) \geq \sum_{\omega} u(a_1, \omega) \mu(\omega|a_2). \quad (2)$$

Then, the expected payoffs of the agent and the regulator under experiment σ are given by:

$$U(\sigma) = \sum_{\omega} \bar{\mu}(\omega) \sum_{k=1,2} \sigma_k(\omega) u(a_k, \omega), \quad (3)$$

$$V(\sigma) = \sum_{\omega} \bar{\mu}(\omega) \sum_{k=1,2} \sigma_k(\omega) v(a_k, \omega). \quad (4)$$

If the agent is unable to conduct an experiment, he chooses an action based on the prior. Let a_p be the agent's prior-optimal action, i.e., $a_p \in \arg \max_a \sum_{\omega} \bar{\mu}(\omega) u(a, \omega)$. The regulator approves an experiment σ if and only if it does not reduce her expected utility relative to the prior-based decision:

$$V(\sigma) \geq \sum_{\omega \in \Omega} \bar{\mu}(\omega) v(a_p, \omega). \quad (5)$$

Otherwise, the regulator blocks the experiment, and the agent chooses action a_p . Therefore, the agent's problem is to select the experiment σ that maximizes his expected utility subject to the regulator's approval constraint:

$$\max_{\sigma \in \Sigma} U(\sigma) \quad \text{s.t.} \quad V(\sigma) \geq \sum_{\omega \in \Omega} \bar{\mu}(\omega) v(a_p, \omega). \quad (6)$$

3 The optimal experiment

Let Σ^* denote the set of experiments that solve the agent's optimization problem in (6), and let $\sigma^* \in \Sigma^*$ be one such optimal experiment. In the absence of a regulator, the agent would choose an experiment that fully reveals the state ω , allowing him to always take his most preferred action $a_g(\omega)$. As usual, we refer to this experiment as the *perfect experiment*, and denote it by σ^{PE} , which is defined as $\sigma^{PE} : \sigma_g(\omega) = 1 \quad \forall \omega \in \Omega$. Thus, if no regulator were present, the agent would implement $\sigma^* = \sigma^{PE}$.

However, because the regulator has the authority to block experiments, the agent may not be able to implement the perfect experiment. The following proposition establishes an important result: if the regulator does not approve the perfect experiment, then any optimal experiment that the agent chooses will leave the regulator with an expected payoff equal to what she would receive if no experiment was carried out. That is, the regulator is never strictly better off compared to the baseline where the agent selects an action based only on the prior.

Proposition 1. *If the perfect experiment σ^{PE} is not approved by the regulator, then for every optimal experiment $\sigma^* \in \Sigma^*$, the regulator's expected payoff remains unchanged compared to the prior, i.e., $V(\sigma^*) = \sum_{\omega \in \Omega} \bar{\mu}(\omega)v(a_p, \omega)$.*

This implies that whenever the agent and the regulator disagree—meaning the perfect experiment is not implementable—the agent must design an experiment that makes the regulator exactly indifferent between accepting and rejecting it. In this case, the approval constraint binds, and any allowed experiment lies on the threshold of what the regulator is willing to accept.

This result reflects the idea that when the agent's interests diverge from those of the regulator, the regulator can prevent harm but cannot enforce improvements. Her role is protective: she ensures that information acquisition does not reduce expected societal welfare, but she cannot compel the agent to generate surplus for others if it conflicts with his own incentives.

To better understand the structure of optimal experiments, we now characterize the set of feasible payoff pairs and the corresponding Pareto frontier. We call a payoff pair $(x, y) \in \mathbb{R}^2$ *feasible* if there exists an incentive-compatible experiment $\sigma \in \Sigma$ such that

$(V(\sigma), U(\sigma)) = (x, y)$. Let F be the set of all such feasible payoff pairs:

$$F := \{(V(\sigma), U(\sigma)) \mid \sigma \in \Sigma : (2)\}.$$

The *Pareto frontier* consists of the maximal elements of F , i.e., those $(x, y) \in F$ for which there exists no $(x', y') \in F$ with $x' \geq x$, $y' \geq y$, and at least one inequality strict.

To describe experiments on the Pareto frontier, we classify states according to whether the agent's and the regulator's preferred actions coincide or not. For each $g, r \in \{1, 2\}$, define:

$$\Omega_{g,r} := \left\{ \omega \in \Omega \mid \arg \max_a u(a, \omega) = a_g, \arg \max_a v(a, \omega) = a_r \right\}.$$

This partitions the state space into four mutually exclusive regions: $\Omega_{1,1}$, $\Omega_{1,2}$, $\Omega_{2,1}$, and $\Omega_{2,2}$. For states in which the agent and the regulator disagree—that is, $\omega \in \Omega_{1,2} \cup \Omega_{2,1}$ —we define the slope of state ω as:

$$\begin{aligned} m(\omega) &:= \frac{v(a_2, \omega) - v(a_1, \omega)}{u(a_1, \omega) - u(a_2, \omega)} \quad \text{for } \omega \in \Omega_{1,2}, \\ m(\omega) &:= \frac{v(a_1, \omega) - v(a_2, \omega)}{u(a_2, \omega) - u(a_1, \omega)} \quad \text{for } \omega \in \Omega_{2,1}. \end{aligned}$$

Finally, we define two sets of states based on how the experiment σ allocates signal probabilities across states where the agent and the regulator disagree:

$$\begin{aligned} \Omega^+(\sigma) &:= \{\omega \in \Omega_{1,2} \mid \sigma_1(\omega) > 0\} \cup \{\omega \in \Omega_{2,1} \mid \sigma_2(\omega) > 0\}, \\ \Omega^-(\sigma) &:= \{\omega \in \Omega_{2,1} \mid \sigma_1(\omega) > 0\} \cup \{\omega \in \Omega_{1,2} \mid \sigma_2(\omega) > 0\}. \end{aligned}$$

Note that a state $\omega \in \Omega$ may appear in both $\Omega^+(\sigma)$ and $\Omega^-(\sigma)$ depending on how the experiment distributes signal probabilities.

The following proposition characterizes the structure of experiments on the Pareto frontier:

Proposition 2. *Let $(U(\sigma), V(\sigma))$ be on the Pareto frontier. Then:*

- a) $\sigma_k(\omega) = 1$ for every $\omega \in \Omega_{k,k}$ and every $k \in \{1, 2\}$;
- b) If $\omega' \in \Omega^-(\sigma)$ and $\omega'' \in \Omega^+(\sigma)$, then $m(\omega') \geq m(\omega'')$.

This result provides two key insights into the structure of optimal experiments when the agent is constrained by the regulator's approval. Part (a) states that in any state where the agent and the regulator agree on the preferred action, the experiment always

reveals the state perfectly. These states do not generate any conflict, so there is no reason to withhold information. Part (b) applies to the states where preferences diverge. It shows that the agent strategically prioritizes states with a higher relative cost to the regulator for full revelation. Specifically, when moving away from the perfect experiment, the agent first distorts information in states with higher slope values $m(\omega)$ (i.e. those in $\Omega^-(\sigma)$) before doing so in states with lower slope values. Intuitively, these are the states where reducing informativeness harms the agent the least relative to the regulator's gain.

Together with Proposition 1, this implies that when the perfect experiment is not acceptable to the regulator, the agent constructs an optimal experiment by selectively deviating from full revelation in the least painful way for regulator. He does so by partially obscuring the states where the relative marginal cost to the regulator (per unit of agent utility) is lowest, continuing until the approval constraint binds. The resulting experiment lies on the Pareto frontier and makes the regulator exactly indifferent between approval and rejection.

Combining Propositions 1 and 2, we build the following search algorithm over the Pareto frontier to characterize the optimal experiment σ^* :

Algorithm 1. *Constructing the Optimal Experiment.*

STEP 1. Start with the perfect experiment σ^{PE} , where $\sigma_g(\omega) = 1$ for all $\omega \in \Omega$.

STEP 2. If the regulator accepts σ^{PE} , implement it. Otherwise, proceed to the next step.

STEP 3. Initialize the experiment by setting $\sigma := \sigma^{PE}$. Identify the state $\omega_1 \in \Omega_{1,2} \cup \Omega_{2,1}$ with the highest slope $m(\omega_1)$ among all such states. If multiple states share the highest slope, select one arbitrarily.

STEP 4. Gradually reduce $\sigma_g(\omega_1)$ from 1 toward 0, keeping all other states at full revelation, until the acceptability constraint binds with equality: $V(\sigma) = \sum_{\omega \in \Omega} \bar{\mu}(\omega)v(a_p, \omega)$. If such a σ is found, implement it. Otherwise, proceed to the next step.

STEP 5. Fix the state with the highest slope $m(\omega_1)$ at $\sigma_g(\omega_1) = 0$, and begin reducing $\sigma_g(\omega_2)$, the second-highest, while keeping all others at full revelation.

STEP 6. Continue this process, each time selecting the next state ω_l with the next-highest slope and reducing $\sigma_g(\omega_l)$, until the regulator accepts.

4 Comparative statics: Multiple regulators

We now extend the analysis to settings in which the agent faces multiple regulators, each of whom must approve the experiment. Such situations arise in a variety of institutional environments, ranging from regulatory appeals processes to committee-based approvals, where decisions are made either sequentially or simultaneously. To analyze these settings, we first introduce a general framework for comparing how different regulators constrain the agent, defining a notion of regulatory strictness and establishing a key comparative result. These foundations will support the applications discussed in Section 5.

Let the agent face n regulators, indexed by the set $N = \{1, \dots, n\}$. Each regulator is a utility maximizer, and we assume that their preferences are *aligned*: that is, they agree on the preferred action in every state. Formally, $a_r^i(\omega) = a_r^j(\omega)$ for all $i, j \in N$ and all $\omega \in \Omega$. Let $\bar{\sigma}_i$ denote the optimal experiment that regulator i would approve if she were acting alone, as in the single-regulator case. We use $\sigma_g^i(\omega)$ (or $\sigma_r^i(\omega)$) to denote the probability that signal a_g (or a_r) is sent in a given state $\omega \in \Omega$ when the agent faces any regulator i .

We say that regulator i *favors the perfect experiment* if she would approve full information revelation when acting alone, that is, if $V_i(\sigma^{PE}) \geq v_i(a_p, \omega)$.

To compare how permissive or restrictive regulators are, we introduce a partial order over regulators based on their preferences. Regulator i is said to be *stricter* than regulator j if the following two conditions hold:

- (i) For all $\omega \in \Omega$, $v_i(a_p, \omega) - v_i(a_g(\omega), \omega) \geq v_j(a_p, \omega) - v_j(a_g(\omega), \omega)$, with strict inequality holding for at least one state.
- (ii) For all $\omega, \omega' \in \Omega_{1,2} \cup \Omega_{2,1}$, if $v_i(a_r(\omega), \omega) - v_i(a_g(\omega), \omega) > v_i(a_r(\omega'), \omega') - v_i(a_g(\omega'), \omega')$, then the same inequality must hold for regulator j .

The first condition captures the idea that a stricter regulator values the prior-based decision more than a more lenient one does, relative to how much the agent values it. In other words, she is more reluctant to allow experiments that might overturn the action of the agent under the prior. The second condition guarantees that states can be compared among regulators. That is, when agent and regulators disagree, the hierarchy among preferred states must be preserved among regulators.

The next result shows that stricter regulators impose tighter constraints on the agent. If regulator i is stricter than regulator j , then the experiment i would approve leads to a higher payoff for j than her own optimal experiment —if she were to face the agent on her own— and yields lower utility for the agent. In simpler words, a more lenient regulator is happier under the constraints of a more stricter regulator than her own.

Proposition 3. *If regulator i is stricter than regulator j , then:*

- (i) $V_j(\bar{\sigma}_i) > V_j(\bar{\sigma}_j)$;
- (ii) $U(\bar{\sigma}_i) < U(\bar{\sigma}_j)$.

The implications of this proposition are straightforward. Since by Proposition 1, we know that the approval constraint must hold with equality, a more lenient regulator will accept some experiments that stricter regulator would reject. These rejected experiments are precisely those that the agent would prefer. This result is key to examine two institutional scenarios in which the agent faces multiple regulators: one where approval can be obtained through a sequence of appeals, and another where all regulators vote simultaneously.

5 Applications

Our model sheds light on different real-world settings in which an agent seeks information before making a decision but faces approval constraints from regulatory institutions. These constraints differ in structure across domains. In some cases, a single regulator evaluates the proposed experiment. In others, the agent may appeal a rejection and re-submit the experiment to another authority. And in many important domains, approval is granted by a committee, either through formal voting rules or consensus procedures.

This section applies the theoretical insights from Section 4 to two broad classes of institutional settings: sequential approval, which models appeals and re-submissions, and collective approval, which captures voting-based mechanisms. For each structure, we relate the core propositions to illustrative examples from clinical testing, warrant investigations, data privacy and medical trials.

5.1 Sequential Approval: Appeals in Clinical Trials and Investigations

Many institutional environments feature sequential oversight, where an agent can escalate a rejected proposal to a higher authority. This occurs, for instance, in pharmaceutical regulation: a clinical trial application rejected by an FDA division may be appealed to a higher-level advisory board or to the Office of Special Medical Programs.²

Similarly, law enforcement officers escalate denied warrant requests to higher judges, or pass an investigation proposal up a departmental hierarchy.³

We model such settings as cases where the agent faces multiple regulators sequentially. Each regulator evaluates the experiment independently. If a regulator approves the proposal, it is immediately implemented. If she rejects, the agent may revise and resubmit to the next regulator in line. This structure maps naturally to our theoretical setup of appeals. Formally, let $i \in N = \{1, \dots, n\}$ denote the order in which the agent approaches the regulators. The following proposition shows that if the regulators are totally ordered by strictness, the final regulator effectively determines the outcome.

Proposition 4. *If regulator i is stricter than regulator j for all $i, j \in N$, then experiment $\bar{\sigma}_n$ will be implemented.*

In particular, this implies that the perfect experiment is implemented if and only if the last regulator in the sequence favors it:

Proposition 5. *The perfect experiment is implemented if and only if the n -th regulator favors the perfect experiment.*

These results highlight the importance of who appears last in the appeals process. The agent benefits from sequential review only if the final regulator is more lenient than the others. Compared to the single-regulator setting, the agent is better off facing multiple sequential regulators when the last one is less strict than the one he would have otherwise faced alone. Conversely, if the final regulator is stricter, the agent ends up worse off.

From a policy perspective, this has clear institutional implications. If a designer (e.g., a government) wants to make it harder for agents to access information, placing the strictest

²See FDA, “Formal Dispute Resolution: Sponsor Appeals Above the Division Level,” *U.S. Food and Drug Administration*, Guidance for Industry, November 2001.

³See Administrative Office of the U.S. Courts (2022), *Guide to Judiciary Policy*, Vol. 4: Probation, Part B: Investigative Policies and Procedures, §270.20 (“Warrant Applications”).

regulator last ensures that no experiment passes unless it meets the highest standard. On the other hand, if the goal is to encourage a freer access to information, the most lenient regulator should be placed at the end of the sequence.

5.2 Collective Approval: Voting in Data Privacy and Medical Review Boards

Other institutional settings involve collective decisions by a panel or committee. Here, regulators evaluate the proposal simultaneously, and approval depends on the aggregation of their judgments—often via majority vote. This structure is common in legal, regulatory, and medical domains. For example, under the GDPR, cross-border data access disputes are resolved by the European Data Protection Board (EDPB), which votes to determine whether corporate practices meet regulatory standards.⁴ Also, Institutional Review Boards (IRBs) and FDA advisory panels use collective voting to approve or deny proposed clinical trials.⁵

We model such environments by letting each regulator vote to accept or reject a proposed experiment. Approval is granted if a majority supports implementation:

$$\frac{1}{N} \sum_{i=1}^N \mathbb{I}(V_i(\sigma) \geq V_i(\bar{\mu})) \geq \frac{1}{2}. \quad (7)$$

As in the sequential case, we assume that regulators' preferences are aligned and ordered by strictness. The following proposition shows that the outcome is determined by the median regulator in this ordering.

Proposition 6. *If the set of regulators N is ordered by strictness, then experiment $\bar{\sigma}_m$ will be implemented, where regulator m is the median regulator in terms of strictness.*

This result highlights the central role of the median regulator under majority rule, resembling the logic of the well-known median voter theorem (Black, 1948). In this setting, the agent is better off under voting than under a single-regulator regime of regulator i if and only if regulator i is stricter than the median regulator m , and worse off otherwise (i.e. if regulator m was stricter than i).

⁴See European Union. (2016). Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 (General Data Protection Regulation), Articles 64–65. Official Journal of the European Union, L 119/1.

⁵See U.S. Department of Health and Human Services. (2022). Protection of Human Subjects, 45 CFR §46.108 (IRB functions and operations).

6 Conclusion

This paper examines how institutional approval constraints reshape the way agents gather and structure information. We develop a model in which an agent can run an experiment to guide his decision, but must first obtain authorization from a regulator. We characterize the structure of optimal experiments and identify the key trade-offs agents face when tailoring information to external approval: If the regulator opposes revelation of full information, the approval constraint forces the agent to distort the experiment in strategic ways, reducing informativeness in states where the regulator is least likely to object. We also extend the framework to environments with multiple regulators, offering comparative predictions across different institutional structures. In sequential settings, the final reviewer determines the outcome; in collective settings, the median regulator is pivotal.

Future work could explore several extensions. One direction is to study dynamic environments. Agents might actually update proposals based on prior approvals, try to build credibility through safe initial experiments, or predict certain approval cycles. Conversely, regulators may adopt history-dependent approval policies, such as “fast-track” rules that reward early success or gradually tighten standards based on observed behavior. These dynamics introduce intertemporal incentives and learning, shifting the problem from one-shot design to repeated strategic interaction. A second direction is to consider that the regulator need not be a fixed, passive evaluator. She might be uncertain, strategic, or internally divided. A review board could reject proposals to signal political toughness or preserve institutional reputation, even if privately indifferent. Committees may aggregate conflicting priorities (e.g., legal caution vs. scientific ambition) leading to approval outcomes shaped by internal bargaining. Modeling the regulator as a strategic or multi-attribute player opens further questions about institutional design and communication.

References

- AGAN, A. AND S. STARR (2018): “Ban the Box, Criminal Records, and Racial Discrimination: A Field Experiment,” *Quarterly Journal of Economics*, 133, 191–235. Cited on pages 1 and 4.
- ALONSO, R. AND O. CÂMARA (2016): “Persuading Voters,” *American Economic Review*, 106, 3590–3605. Cited on page 5.
- BARDHI, A. AND Y. GUO (2018): “Modes of persuasion toward unanimous consent,” *Theoretical Economics*, 13, 1111–1149. Cited on page 5.
- BLACK, D. (1948): “On the Rationale of Group Decision-making,” *Journal of Political Economy*, 56, 23–34. Cited on page 13.
- DOVAL, L. AND J. C. ELY (2020): “Sequential information design,” *Econometrica*, 88, 2575–2608. Cited on page 5.
- DOVAL, L. AND V. SKRETA (2024): “Constrained Information Design,” *Mathematics of Operations Research*, 49, 78–106. Cited on page 4.
- GALPERTI, S. AND J. PEREGO (2023): “Privacy and the Value of Data,” *AEA Papers and Proceedings*, 113, 197–203. Cited on page 1.
- HACKMANN, M. B., J. T. KOLSTAD, AND A. E. KOWALSKI (2015): “Adverse Selection and an Individual Mandate: When Theory Meets Practice,” *American Economic Review*, 105, 1030–1066. Cited on page 1.
- HENRY, E. AND M. OTTAVIANI (2019): “Research and the Approval Process: The Organization of Persuasion,” *American Economic Review*, 109, 911–955. Cited on page 4.
- HIRSHLEIFER, J. (1971): “The private and social value of information and the reward to inventive activity,” *American Economic Review*, 61, 561–574. Cited on pages 1 and 4.
- HU, P. AND J. SOBEL (2022): “Getting Permission,” *American Economic Review: Insights*, 4, 459–472. Cited on page 5.

- ICHIHASHI, S. (2019): “Limiting Sender’s Information in Bayesian Persuasion,” *Games and Economic Behavior*, 117, 276–288. Cited on page 4.
- (2020): “Online privacy and information disclosure by consumers,” *American Economic Review*, 110, 569–595. Cited on page 4.
- KAMENICA, E. AND M. GENTZKOW (2011): “Bayesian Persuasion,” *American Economic Review*, 101, 2590–2615. Cited on page 4.
- MCCLELLAN, A. (2022): “Experimentation and Approval Mechanisms,” *Econometrica*, 90, 2215–2247. Cited on page 4.
- MORRIS, S. AND H. S. SHIN (2002): “Social value of public information,” *American Economic Review*, 92, 1521–1534. Cited on page 4.
- ONUCHIC, P. AND D. RAY (2023): “Signaling and discrimination in collaborative projects,” *American Economic Review*, 113, 210–252. Cited on page 1.
- STRACK, P. AND K. H. YANG (2024): “Privacy-Preserving Signals,” *Econometrica*, 92, 1907–1938. Cited on pages 1 and 4.
- TIROLE, J. (2016): “From bottom of the barrel to cream of the crop: Sequential screening with positive selection,” *Econometrica*, 84, 1291–1343. Cited on page 4.
- WITTBRODT, C. AND N. YODER (2025): “Delegated Experimentation and Approval,” Working paper. Cited on page 5.

Appendix

A Proofs

A.1 Proof of Proposition 1

Proof. Let σ^{PE} be the perfect experiment and let σ^* be an optimal experiment. If σ^{PE} is not approved, then $\sigma^{PE} \neq \sigma^*$, so that there exists a state $\omega' \in \Omega$ with $\bar{\mu}(\omega') > 0$ and $\sigma_g^*(\omega') < 1$, where $a_g(\omega') = a_g$. It follows that $U(\sigma^*) < U(\sigma^{PE})$. For $\varepsilon \in [0, 1]$, define the experiment $\sigma^\varepsilon \in \Sigma$ by $\sigma_1^\varepsilon(\omega) = \varepsilon\sigma_1^{PE}(\omega) + (1-\varepsilon)\sigma_1^*(\omega)$ and $\sigma_2^\varepsilon(\omega) = \varepsilon\sigma_2^{PE}(\omega) + (1-\varepsilon)\sigma_2^*(\omega)$ for every $\omega \in \Omega$. Because both σ^{PE} and σ^* are incentive compatible, so is σ^ε , and

$$\begin{aligned} U(\sigma^\varepsilon) &= \varepsilon U(\sigma^{PE}) + (1-\varepsilon)U(\sigma^*), \\ V(\sigma^\varepsilon) &= \varepsilon V(\sigma^{PE}) + (1-\varepsilon)V(\sigma^*). \end{aligned}$$

Since $U(\sigma^*) < U(\sigma^{PE})$, we have $U(\sigma^\varepsilon) > U(\sigma^*)$ for every $\varepsilon > 0$, and if $V(\sigma^*) < \sum_{\omega \in \Omega} \bar{\mu}(\omega)v(a_p, \omega)$, then $V(\sigma^\varepsilon) < \sum_{\omega \in \Omega} \bar{\mu}(\omega)v(a_p, \omega)$ for small $\varepsilon > 0$. By the optimality of σ^* , it follows that $V(\sigma^*) = \sum_{\omega \in \Omega} \bar{\mu}(\omega)v(a_p, \omega)$. $\square$

A.2 Proof of Proposition 2

Proof. Part a) We consider here the case that $k = 1$; the other case is analogous and therefore omitted. By contradiction, suppose there exists $\omega' \in \Omega_{11}$ with $\sigma_1(\omega') < 1$. Consider $\tilde{\sigma}$ with

$$\tilde{\sigma}_1(\omega) = \begin{cases} \sigma_1(\omega) & \text{if } \omega \neq \omega', \\ 1 & \text{if } \omega = \omega'. \end{cases}$$

We show that $\tilde{\sigma}$ is incentive compatible. Thus, we need to show

$$\begin{aligned} \sum_{\omega} (u(a_1, \omega) - u(a_2, \omega))\tilde{\sigma}_1(\omega)\bar{\mu}(\omega) &\geq 0, \\ \sum_{\omega} (u(a_2, \omega) - u(a_1, \omega))\tilde{\sigma}_2(\omega)\bar{\mu}(\omega) &\geq 0. \end{aligned}$$

Because $\omega' \in \Omega_{11}$, we have $u(a_1, \omega') - u(a_2, \omega') > 0$. Hence, since σ is incentive compatible,

$$\begin{aligned} \sum_{\omega} (u(a_1, \omega) - u(a_2, \omega))\tilde{\sigma}_1(\omega)\bar{\mu}(\omega) &= \sum_{\omega \neq \omega'} (u(a_1, \omega) - u(a_2, \omega))\sigma_1(\omega)\bar{\mu}(\omega) + (u(a_1, \omega') - u(a_2, \omega'))\bar{\mu}(\omega) \\ &> \sum_{\omega} (u(a_1, \omega) - u(a_2, \omega))\sigma_1(\omega)\bar{\mu}(\omega) \\ &\geq 0 \end{aligned}$$

and

$$\begin{aligned}
\sum_{\omega} (u(a_2, \omega) - u(a_1, \omega)) \tilde{\sigma}_2(\omega) \bar{\mu}(\omega) &= \sum_{\omega \neq \omega'} (u(a_2, \omega) - u(a_1, \omega)) \sigma_2(\omega) \bar{\mu}(\omega) \\
&> \sum_{\omega} (u(a_2, \omega) - u(a_1, \omega)) \sigma_2(\omega) \bar{\mu}(\omega) \\
&\geq 0.
\end{aligned}$$

Thus, $(U(\tilde{\sigma}), V(\tilde{\sigma})) \in F$.

Now, since $\omega' \in \Omega_{11}$, we have

$$\begin{aligned}
U(\tilde{\sigma}) - U(\sigma) &= \bar{\mu}(\omega')(1 - \sigma_1(\omega'))(u(a_1, \omega') - u(a_2, \omega')) > 0, \\
V(\tilde{\sigma}) - V(\sigma) &= \bar{\mu}(\omega')(1 - \sigma_1(\omega'))(v(a_1, \omega') - v(a_2, \omega')) > 0;
\end{aligned}$$

a contradiction to $(U(\sigma), V(\sigma))$ being in the Pareto frontier. Thus, $\sigma_1(\omega) = 1$ for every $\omega \in \Omega_{11}$.

Part b) By contradiction, suppose there exist $\omega' \in \Omega^-(\sigma)$ and $\omega'' \in \Omega^+(\sigma)$ with $m(\omega') < m(\omega'')$. We consider here the case that $\omega' \in \Omega_{12}$ and $\omega'' \in \Omega_{21}$; the other cases are analogous and therefore omitted. Since $\sigma_1(\omega'), \sigma_1(\omega'') < 1$, we can find $\varepsilon', \varepsilon'' > 0$ such that

$$\sigma_1(\omega') + \varepsilon', \sigma_1(\omega'') + \varepsilon'' < 1$$

and

$$\varepsilon' \bar{\mu}(\omega')(u(a_1, \omega') - u(a_2, \omega')) = \varepsilon'' \bar{\mu}(\omega'')(u(a_2, \omega'') - u(a_1, \omega'')). \quad (8)$$

Observe that (8) implies

$$\varepsilon' \bar{\mu}(\omega') m(\omega')(u(a_1, \omega') - u(a_2, \omega')) < \varepsilon'' \bar{\mu}(\omega'') m(\omega'')(u(a_2, \omega'') - u(a_1, \omega'')). \quad (9)$$

Consider $\tilde{\sigma}$ with

$$\tilde{\sigma}_1(\omega) = \begin{cases} \sigma_1(\omega) & \text{if } \omega \neq \omega', \omega'', \\ \sigma_1(\omega) + \varepsilon' & \text{if } \omega = \omega', \\ \sigma_1(\omega) + \varepsilon'' & \text{if } \omega = \omega''. \end{cases}$$

We show that $\tilde{\sigma}$ is incentive compatible. Since σ is incentive compatible, (8) implies

$$\begin{aligned}
\sum_{\omega} \tilde{\sigma}_1(\omega) \bar{\mu}(\omega) (u(a_1, \omega) - u(a_2, \omega)) &= \sum_{\omega} \sigma_1(\omega) \bar{\mu}(\omega) (u(a_1, \omega) - u(a_2, \omega)) \\
&\quad + \varepsilon' \bar{\mu}(\omega') (u(a_1, \omega') - u(a_2, \omega')) + \varepsilon'' \bar{\mu}(\omega'') (u(a_1, \omega'') - u(a_2, \omega'')) \\
&= \sum_{\omega} \sigma_1(\omega) \bar{\mu}(\omega) (u(a_1, \omega) - u(a_2, \omega)) \\
&\geq 0
\end{aligned}$$

and

$$\begin{aligned}
\sum_{\omega} \tilde{\sigma}_2(\omega) \bar{\mu}(\omega) (u(a_2, \omega) - u(a_1, \omega)) &= \sum_{\omega} \sigma_2(\omega) \bar{\mu}(\omega) (u(a_2, \omega) - u(a_1, \omega)) \\
&\quad - \varepsilon' \bar{\mu}(\omega') (u(a_2, \omega') - u(a_1, \omega')) - \varepsilon'' \bar{\mu}(\omega'') (u(a_2, \omega'') - u(a_1, \omega'')) \\
&= \sum_{\omega} \sigma_2(\omega) \bar{\mu}(\omega) (u(a_2, \omega) - u(a_1, \omega)) \\
&\geq 0.
\end{aligned}$$

Thus, $(U(\tilde{\sigma}), V(\tilde{\sigma})) \in F$.

Lastly, we compare the expected payoffs generated by $\tilde{\sigma}$ and σ , respectively. By (8),

$$U(\tilde{\sigma}) - U(\sigma) = \varepsilon' \bar{\mu}(\omega') (u(a_1, \omega') - u(a_2, \omega')) + \varepsilon'' \bar{\mu}(\omega'') (u(a_1, \omega'') - u(a_2, \omega'')) = 0;$$

by (9),

$$\begin{aligned}
V(\tilde{\sigma}) - V(\sigma) &= \varepsilon' \bar{\mu}(\omega') (v(a_1, \omega') - v(a_2, \omega')) + \varepsilon'' \bar{\mu}(\omega'') (v(a_1, \omega'') - v(a_2, \omega'')) \\
&= \varepsilon' \bar{\mu}(\omega') m(\omega') (u(a_2, \omega') - u(a_1, \omega')) + \varepsilon'' \bar{\mu}(\omega'') m(\omega'') (u(a_2, \omega'') - u(a_1, \omega'')) \\
&> 0.
\end{aligned}$$

This is a contradiction to $(U(\sigma), V(\sigma))$ being in the Pareto frontier. Thus, $m(\omega') \geq m(\omega'')$. $\square$

A.3 Proof of Proposition 3

Proof. Part i). Let regulator i be stricter than regulator j . We need to show that $V_j(\bar{\sigma}_i) > V_j(\bar{\sigma}_j), \forall i, j \in N$. First, rearranging $V_j(\bar{\sigma}_i)$ yields:

$$\begin{aligned}
V_j(\bar{\sigma}_i) &= \sum_{\omega} \bar{\mu}(\omega) \sum_{g=1,2} \bar{\sigma}_g^i(\omega) v_j(a_g, \omega) \\
&= \sum_{\omega \in \Omega} \bar{\mu}(\omega) [\bar{\sigma}_r^i(\omega) v_j(a_r, \omega) + \bar{\sigma}_g^i(\omega) v_j(a_g, \omega)] \\
&= \sum_{\omega \in \Omega} \bar{\mu}(\omega) [\bar{\sigma}_r^i(\omega) (v_j(a_r, \omega) - v_j(a_g, \omega)) + v_j(a_g, \omega)] \tag{10}
\end{aligned}$$

We have to show that:

$$V_j(\bar{\sigma}_i) - V_j(\bar{\sigma}_j) > 0;$$

Plugging equation (10) yields:

$$\begin{aligned}
V_j(\bar{\sigma}_i) - V_j(\bar{\sigma}_j) &= \sum_{\omega \in \Omega} \bar{\mu}(\omega) [\bar{\sigma}_r^i(\omega) (v_j(a_r, \omega) - v_j(a_g, \omega))] - \sum_{\omega \in \Omega} \bar{\mu}(\omega) [\bar{\sigma}_r^j(\omega) (v_j(a_r, \omega) - v_j(a_g, \omega))] \\
&\tag{11}
\end{aligned}$$

$$= \sum_{\omega \in \Omega} \bar{\mu}(\omega) [(\bar{\sigma}_r^i(\omega) - \bar{\sigma}_r^j(\omega)) (v_j(a_r, \omega) - v_j(a_g, \omega))] > 0 \tag{12}$$

Note that, since regulator preferences are aligned, $\sigma_r^i(\omega) = \sigma_r^j(\omega) = 1$ for all $\omega \in \Omega_{1,1} \cup \Omega_{2,2}$. Thus it is enough to show that $\sigma_r^i(\omega) \geq \sigma_r^j(\omega)$ for all $\omega \in \Omega_{1,2} \cup \Omega_{2,1}$ to prove the proposition.

Because i is the stricter regulator, set $\bar{\sigma}_i \neq \sigma^{PE}$ without loss of generality. Now:

Case 1: $\bar{\sigma}_j = \sigma^{PE}$

Since $\sigma^{PE}(\omega) = \sigma_g^j = 1$ for all $\omega \in \Omega$, and $\bar{\sigma}_i \neq \sigma^{PE}$, then $\bar{\sigma}_g^i < 1$ for at least one $\omega \in \Omega_{1,2} \cup \Omega_{2,1}$. Since the set $\Omega_{1,2} \cup \Omega_{2,1}$ is the same for all regulators, $\sigma_r^i(\omega) \geq \sigma_r^j(\omega)$ for all $\omega \in \Omega_{1,2} \cup \Omega_{2,1}$

Case 2: $\bar{\sigma}_j \neq \sigma^{PE}$

Since $\bar{\sigma}_j \neq \sigma^{PE}$, we know that by Proposition 1, $V_i(\bar{\sigma}_i) = V_i(\bar{\mu})$. Rearranging yields:

$$V_i(\bar{\sigma}_i) = \sum_{\omega \in \Omega} \bar{\mu}(\omega) v_i(a_p, \omega) \tag{13}$$

$$\sum_{\omega \in \Omega} \bar{\mu}(\omega) \bar{\sigma}_r^i(\omega) (v_i(a_r, \omega) - v_i(a_g, \omega)) + v_i(a_g, \omega) = \sum_{\omega \in \Omega} \bar{\mu}(\omega) (v_i(a_g, \omega) - v_i(a_p, \omega)) \tag{14}$$

Now partition Ω into $\Omega_{1,1} \cup \Omega_{2,2}$ and $\Omega_{1,2} \cup \Omega_{2,1}$ and note that $\sigma_r^i(\omega) = \sigma_r^j(\omega) = 1$ for all $\omega \in \Omega_{1,1} \cup \Omega_{2,2}$, and further note that $\sum_{\omega \in \Omega_{1,2} \cup \Omega_{2,1}} \bar{\mu}(v_i(a_p, \omega) - v_i(a_g, \omega)) = \sum_{\omega \in \Omega_{1,2} \cup \Omega_{2,1}} \bar{\mu}(v_i(a_r, \omega) - v_i(a_g, \omega))$ since either $a_p = a_r$ or $a_p = a_g$ for all $\omega \in \Omega_{1,2} \cup \Omega_{2,1}$ and $\sum_{\omega \in \Omega_{1,2} \cup \Omega_{2,1}} \bar{\mu}(v_i(a_p, \omega) - v_i(a_g, \omega)) = 0$ for states in the latter case.

Rearranging equation (14) yields:

$$\begin{aligned} \sum_{\omega \in \Omega_{1,2} \cup \Omega_{2,1}} \bar{\mu}(\omega) \bar{\sigma}_r^i(\omega) (v_i(a_r, \omega) - v_i(a_g, \omega)) &= \sum_{\omega \in \Omega_{1,2} \cup \Omega_{2,1}} \bar{\mu}(\omega) (v_i(a_r, \omega) - v_i(a_g, \omega)) \\ &\quad - \sum_{\omega \in \Omega_{1,1} \cup \Omega_{2,2}} \bar{\mu}(\omega) (v_i(a_g, \omega) - v_i(a_p, \omega)) \end{aligned} \quad (15)$$

Now, note that applying Proposition 2 and Algorithm 1, one can partition the set $\Omega_{1,2} \cup \Omega_{2,1}$ into three sets $\{\Omega_0, \Omega_1, \text{ and } \Omega_{0-1}\}$ defined as follows:

- Ω_0 is the set of all $\omega \in \Omega_{1,2} \cup \Omega_{2,1}$ such that $\bar{\sigma}_r^i(\omega) = 0$.
- Ω_1 is the set of all $\omega \in \Omega_{1,2} \cup \Omega_{2,1}$ such that $\bar{\sigma}_r^i(\omega) = 1$.
- Ω_{0-1} is a singleton in $\Omega_{1,2} \cup \Omega_{2,1}$ for which $\bar{\sigma}_r^i(\omega) \in (0, 1)$. Let this state be called $\tilde{\omega}$.

Applying the above-mentioned partition to (15) yields:

$$\begin{aligned} \bar{\mu}(\tilde{\omega}) \bar{\sigma}_r^i(\tilde{\omega}) (v_i(a_r, \tilde{\omega}) - v_i(a_g, \tilde{\omega})) &= \sum_{\omega \in \Omega_0 \cup \Omega_{0-1}} \bar{\mu}(\omega) (v_i(a_r, \omega) - v_i(a_g, \omega)) \\ &\quad - \sum_{\omega \in \Omega_{1,1} \cup \Omega_{2,2}} \bar{\mu}(\omega) (v_i(a_g, \omega) - v_i(a_p, \omega)) \end{aligned} \quad (16)$$

Solving for $\bar{\sigma}_r^i(\tilde{\omega})$ yields:

$$\bar{\sigma}_r^i(\tilde{\omega}) = 1 + \frac{\sum_{\omega \in \Omega_0} \bar{\mu}(\omega) (v_i(a_r, \omega) - v_i(a_g, \omega)) - \sum_{\omega \in \Omega_{1,1} \cup \Omega_{2,2}} \bar{\mu}(\omega) (v_i(a_g, \omega) - v_i(a_p, \omega))}{\bar{\mu}(\tilde{\omega}) (v_i(a_r, \tilde{\omega}) - v_i(a_g, \tilde{\omega}))} \quad (17)$$

Applying the same procedure for $V_j(\bar{\sigma}_j)$ yields:

$$\bar{\sigma}_r^j(\tilde{\omega}) = 1 + \frac{\sum_{\omega \in \Omega_0} \bar{\mu}(\omega) (v_j(a_r, \omega) - v_j(a_g, \omega)) - \sum_{\omega \in \Omega_{1,1} \cup \Omega_{2,2}} \bar{\mu}(\omega) (v_j(a_g, \omega) - v_j(a_p, \omega))}{\bar{\mu}(\tilde{\omega}) (v_j(a_r, \tilde{\omega}) - v_j(a_g, \tilde{\omega}))} \quad (18)$$

Now, first note that trivially $\bar{\sigma}_r^i(\omega) = \bar{\sigma}_r^j(\omega) = 0$ for all $\omega \in \Omega_0$ and $\bar{\sigma}_r^i(\omega) = 1 \geq \bar{\sigma}_r^j$ for all $\omega \in \Omega_1$. Finally, for state $\tilde{\omega}$, we know that since the second term of the RHS of equations (17) and (18) must be negative, and given that regulator i is stricter than regulator j , it directly follows that $\bar{\sigma}_r^i(\tilde{\omega}) > \bar{\sigma}_r^j(\tilde{\omega})$.

Therefore, $\sigma_r^i(\omega) \geq \sigma_r^j(\omega)$ for all $\omega \in \Omega_{1,2} \cup \Omega_{2,1}$, where at least one inequality holds strictly. Thus $V_j(\bar{\sigma}_i) > V_j(\bar{\sigma}_j)$

Part ii) The proof is analogous to Part i). We need to show that $U(\bar{\sigma}_j) - U(\bar{\sigma}_i) > 0$. Rearranging this expression by following the same steps as Part i) yields:

$$U(\bar{\sigma}_j) - U(\bar{\sigma}_i) = \sum_{\omega \in \Omega} \bar{\mu}(\omega) [(\bar{\sigma}_r^i(\omega) - \bar{\sigma}_r^j(\omega))(u(a_g, \omega) - u(a_r, \omega))] > 0 \quad (19)$$

By definition $u(a_g, \omega) \geq u(a_r, \omega)$ for all $\omega \in \Omega$. Additionally, from Part i) we know that $\sigma_r^i(\omega) \geq \sigma_r^j(\omega)$ for all $\omega \in \Omega_{1,2} \cup \Omega_{2,1}$, thus $U(\bar{\sigma}_j) > U(\bar{\sigma}_i)$ since at least one of the previous inequalities must hold strictly. $\square$

A.4 Proof of Proposition 4

Proof. Assume that for all $i, j \in N$, regulator i is stricter than regulator j . Note that if the agent reaches decision node n , then experiment $\bar{\sigma}_n$ will be implemented.

First, suppose that an arbitrary regulator i is stricter than regulator n . Then, by Proposition 3 (ii) we have $U(\bar{\sigma}_n) > U(\bar{\sigma}_i)$. Thus, the agent can always offer $\bar{\sigma}_n$ to any regulator i that is stricter than regulator n . Although by Proposition 3 (i) we have $V_i(\bar{\sigma}_n) < V_i(\bar{\sigma}_i)$, regulator i has no incentive to reject $\bar{\sigma}_n$, since he knows that at decision node n experiment $\bar{\sigma}_n$ will be implemented.

Now, suppose instead that regulator i is less strict than regulator n . Then, $U(\bar{\sigma}_i) > U(\bar{\sigma}_n)$. However, if $\bar{\sigma}_i$ were offered to regulator i , he would always reject it because $V_i(\bar{\sigma}_n) > V_i(\bar{\sigma}_i)$, and regulator i knows that at decision node n , experiment $\bar{\sigma}_n$ would be implemented. Consequently, the agent does not have an incentive to deviate from offering $\bar{\sigma}_n$, and therefore the experiment is implemented. $\square$

A.5 Proof of Proposition 5

Proof. ($\Rightarrow$) Let the n -th regulator favor the Perfect Experiment (i.e. $V_n(\sigma^{PE}) - v_n(\bar{\mu}) \geq 0$). Since σ^{PE} is the agent's most preferred experiment, the agent will offer σ^{PE} to the n -th

regulator, which will be accepted. Since at decision node n , σ^{PE} is guaranteed to be accepted, there is no incentive for any regulator $i \in N$ to reject σ^{PE} , and the perfect experiment is implemented.

($\Leftarrow$) Let the n -th regulator not favor the Perfect Experiment (i.e. $V_n(\sigma^{PE}) - v_n(\bar{\mu}) < 0$). Since σ^{PE} will be rejected at decision node n , the agent will offer regulator n the optimal experiment under her participation constraint. Let such an experiment be denoted as $\bar{\sigma}_n$, and note that $V_n(\bar{\sigma}_n) = V_n(\bar{\mu})$, and $U(\bar{\sigma}_n), V_n(\bar{\sigma}_n)$ is in the Pareto frontier of regulator n . This implies that $V_n(\bar{\sigma}_n) > V_n(\sigma^{PE})$ and $\bar{\sigma}_n(\omega) \neq \sigma^{PE}(\omega)$ for at least one $\omega \in \Omega_{1,2} \cup \Omega_{2,1}$. Since all regulator's preferences are aligned, the set $\Omega_{1,2} \cup \Omega_{2,1}$ is the same for all regulators $i \in N$. This means that regulator n is stricter than any regulator i who does not favor σ^{PE} . Thus, by Proposition 3 (Case 1), $V_i(\bar{\sigma}_n) > V_i(\sigma^{PE}), \forall i \in N$, implying that the Perfect Experiment will never be implemented since at least experiment $\bar{\sigma}_n$ is weakly preferred by all regulators. $\square$

A.6 Proof of Proposition 6

Proof. Without loss of generality, suppose the regulators are indexed so that for each $i = 1, \dots, n$, regulator $i + 1$ is strictly stricter than regulator i . Define the median regulator m as follows:

- If N is odd, then $m = \frac{N+1}{2}$.
- If N is even, then $m = \frac{N}{2}$ (so that the stricter side has one more regulator).

Further note that by Proposition 3, the agent will always prefer to implement experiments that leave less strict regulators indifferent, i.e.,

$$U(\bar{\sigma}_i) > U(\bar{\sigma}_j) \quad \forall i < j,$$

while any regulator i will prefer implementing experiments that leave stricter regulators indifferent, i.e.,

$$V_i(\bar{\sigma}_i) < V_i(\bar{\sigma}_j) \quad \forall i < j.$$

Now, for the sake of contradiction, suppose that the agent proposes an experiment $\sigma' \neq \bar{\sigma}_m$. We consider two cases:

Case 1: Suppose that $\sigma' = \bar{\sigma}_i$ for some $i > m$. Then, by the preference ordering for the agent (Proposition 3),

$$U(\bar{\sigma}_i) < U(\bar{\sigma}_m).$$

Hence, the agent strictly prefers offering $\bar{\sigma}_m$ to offering $\bar{\sigma}_i$. This contradicts the assumption that the agent would deviate to σ' .

Case 2: Suppose that $\sigma' = \bar{\sigma}_i$ for some $i < m$. In this case, by Proposition 3, for any regulator j who is stricter than regulator m , we have

$$V_j(\bar{\sigma}_i) < V_j(\bar{\sigma}_j) = V_j(\bar{\mu}).$$

Thus, each such regulator j would reject $\bar{\sigma}_i$. Since the set of regulators who are stricter than regulator m constitutes more than half of the regulators (by the definition of the median), it follows that

$$\frac{1}{N} \sum_{j=1}^N \mathbb{I}(V_j(\bar{\sigma}_i) \geq V_j(\bar{\mu})) < \frac{1}{2}.$$

Hence, the proposal $\bar{\sigma}_i$ would fail to secure simple majority.

Since in both cases any deviation from offering $\bar{\sigma}_m$ either results in a lower payoff for the agent or fails to be accepted by a majority of regulators, the agent has no incentive to deviate from offering $\bar{\sigma}_m$. Consequently, experiment $\bar{\sigma}_m$ will be implemented. $\square$