

Sequential Mechanisms for Evidence Acquisition¹

Elchanan Ben-Porath²

Eddie Dekel³

Barton L. Lipman⁴

First Draft

June 2023

Current Draft

August 2026

¹We thank the National Science Foundation, grant SES–1919319 (Dekel and Lipman), and the US–Israel Binational Science Foundation for support for this research. We thank Simon Board and Duarte Goncalves for pushing us to consider the relationship of Weitzman (1979) to our results and Axel Niemeyer, the editor, and the referees for useful comments and suggestions.

²Department of Economics and Center for Rationality, Hebrew University. Email: benporat@math.huji.ac.il.

³Economics Department, Northwestern University, and School of Economics, Tel Aviv University. Email: eddiedekel@gmail.com.

⁴Department of Economics, Boston University. Email: blipman@bu.edu.

Abstract

We consider optimal mechanisms for inducing agents to acquire costly evidence in a setting where a principal has a good to allocate that all agents want. We show that optimal mechanisms are sequential and have a threshold structure. Agents with higher costs of obtaining evidence and/or worse distributions of value for the principal are asked for evidence later, if at all. We derive these results in part by exploiting the relationship between the Lagrangian for this problem and the classic Weitzman (1979) “Pandora’s box” problem together with Border’s (1991) characterization of feasible interim allocations.

1 Introduction

A principal has a single unit of a good or resource to allocate to one of N agents under uncertainty regarding the value he would receive from allocating it to any given one of them. Each agent wants the good, independently of the value her receiving it provides the principal. Each agent can obtain information which would reveal to her the value she would provide the principal if she receives the good and provide evidence proving this value to the principal. However, this information is costly to the agent, so she would not be willing to get it without a high enough chance that this will lead to her receiving the good. In situations where the number of agents is large and the cost of demonstrating one's credentials is reasonably high relative to the value of receiving the good, the principal faces a difficult incentive problem.

As examples of this situation, consider the head of an organization with multiple divisions, such as a university with multiple departments. The head of the organization has discrete resources to allocate, such as prestigious projects or assignments, or, in the case of a university, job slots. Some divisions would use this resource in ways that are more productive for the organization than others, but all divisions prefer to receive it, independently of their productivity. For a division to determine what value it would produce for the organization if it receives the resource is costly in time and/or effort, so that the division may prefer not to make a serious proposal to receive the resource if it is unlikely to succeed.

As other examples, consider an agency allocating grants for research or contracts to undertake a project. The applicant could invest more or less in preparing their application or bid, by, for example, carrying out research in the former case or developing a prototype in the latter. This would enable the applicant to learn something they can then prove about their ability to carry out the research or see the project to successful completion. However, if the applicant's chance of success is too low, it will not be worthwhile to invest. Similarly, a job applicant might invest more or less in preparing for an interview or other test, but has less incentive to do so if the chances of success are low.

We assume monetary transfers cannot be used. Intuitively, divisions of an organization have funds available to them to carry out actions on behalf of the organization. So it would be counterproductive to the organization to have divisions "bid" for these resources. Furthermore, the willingness to pay may be higher for divisions that are less valuable for the organization to choose, making it not obvious that such transfers could be fruitfully employed.

We characterize optimal mechanisms for the principal in this setting. In doing so, we assume the principal has a high degree of commitment power, to an extent not necessarily natural for some of the examples above. Similarly, we assume that agents have no

private information ex ante, again a more plausible assumption in some settings than others. Where these assumptions are less natural, our results are best interpreted as a benchmark, establishing the best a principal could hope to achieve.

First, assume agents are symmetric both in their costs of obtaining evidence and in the probability distribution over the value they would provide the principal if allocated the resource. Let $c \in (0, 1)$ denote the cost, let v_i denote a typical realization of the value to the principal of giving the good to agent i where these variables are iid across agents, and normalize the value of receiving the good to an agent to 1.

A natural mechanism to consider is for the principal to ask all agents to provide evidence, awarding the good to the agent who proves the highest value. If the number of agents, N , is large, this will not induce all agents to seek evidence. Specifically, if $c > 1/N$, the cost exceeds the expected benefit.

The principal could exclude some agents from the mechanism and only ask some smaller number, say n , such that $1/n \geq c$. As we show later, in a symmetric setting like this, it is *never* optimal for the principal to exclude some agents. Intuitively, the reason is that the principal is better off using a sequential mechanism where he can tailor the number of agents he asks for evidence to the results he sees, asking a larger number of agents when he sees lower values and fewer when he sees high ones. Each agent will still have a $1/N$ chance of receiving the good, but will have lower expected costs because they are not always asked for evidence.

In this symmetric problem, the optimal mechanism is easy to describe. The principal chooses a random ordering of the agents where all orderings are equally likely. He goes to the first agent in the selected ordering and asks her to provide evidence. In equilibrium, she will pay the cost c , learn her value, and prove this value to the principal. If her value is above a certain threshold, v^* , he stops and gives her the good. Otherwise, he continues to the next agent, applying the same threshold to decide whether to give her the good or continue. If all agents have values below v^* , he will end up asking all of them for evidence. In this case, he gives the good to the agent with the highest value.

An important point is that when the principal asks an agent for evidence, he does not tell her where she is in the sequence. To see why this is valuable, consider for simplicity the case of two agents. Suppose that when the agent is asked for evidence, she knows that she is second in line. Then she knows that the other agent's value is below the threshold. Letting F denote the common cdf for v , her probability of getting the good is then $1 - F(v^*) + (1/2)F(v^*) > 1/2$. This is because she certainly gets it if her value is above v^* and she is symmetric to the other agent conditional on her value being below v^* and so gets the good with probability $1/2$ in this event. On the other hand, if she knows she is the first one to be asked, then her probability of getting the good must be strictly smaller. More specifically, it is $1 - F(v^*) + (1/2)[F(v^*)]^2$. This is because she knows

that if she's below the threshold, the second agent will be asked for evidence and she'll only have a chance of getting the good if the other agent is also below the threshold. In short, the second agent has a larger incentive to pay for evidence than the first. Since the agents are symmetric, it is optimal to equalize the incentives by randomizing 50–50 over which agent is asked first, rather than to distort the allocation differently for the two.¹

In some settings, it may be difficult for the principal to ensure that agents learn nothing about their position in the ordering. If agents know how long ago the process started, they are likely to have an estimate of how many agents were asked before them. Like our assumption that the principal can commit to a mechanism, in such settings, our analysis is best thought of as providing a benchmark as a natural first step for analysis.

When the agents are asymmetric, new considerations arise. Because the principal has to give the good with enough probability to any agent he asks for evidence, agents with higher probabilities on high values are better to ask earlier. Because agents with higher costs must be given more incentive to induce them to obtain evidence, he may wish to ask them later.

The optimal mechanism changes in two ways. First, instead of comparing the agent's *values* to the threshold or to one another, we compare *virtual values*. Specifically, for each agent i , instead of comparing v_i to the threshold or another agent's value, we compare $v_i + \lambda_i$ where λ_i reflects the severity of i 's incentive compatibility constraint. In fact, λ_i is the Lagrange multiplier on this constraint: agents who are harder to incentivize are given a constant “advantage” in the form of “extra points” added to their values.

Second, not surprisingly, the randomization over orders is no longer uniform. Indeed, in some cases, the asymmetries across agents will lead to some or all aspects of the ordering being deterministic. Intuitively, if one agent has much lower costs than another, then that agent can more easily bear the burden of being asked for evidence first.

The most general ordering of an optimal mechanism involves what we refer to as *tiers*. Specifically, the agents are partitioned into tiers.² The principal starts with the highest tier, asking agents in some random order for evidence. In each case, he compares the agent's virtual value to the tier 1 threshold, stopping and giving the good to the agent if her virtual value is above the threshold, continuing otherwise. If all the tier 1 agents have virtual values below the tier 1 threshold, the principal will learn all of their values. At this point, he *lowers* the threshold. If any of the agents have virtual values above

¹More specifically, it is not hard to see that there must be some $\bar{v} > v^*$ where the incentive constraints hold with equality if the principal randomizes 50–50 over which agent to start with and uses threshold $\bar{v}$. Because $\bar{v} > v^*$, this mechanism gives the principal a higher expected payoff.

²As we discuss in more detail in Section 3, agents that are “similar enough” will be assigned to the same tier. Hence two agents being placed in the same tier is not a “non-generic” feature.

this lower threshold, he gives the good to the one with the highest virtual value. If not, he continues to the tier 2 agents now using this lower threshold for the tier 2 threshold and continues in this fashion to lower tiers as needed. If none of the agents has a virtual value above the lowest threshold, he asks all for evidence and gives the good to the agent with the highest virtual value. This structure has the uniform randomization with all agents in the same tier as the most symmetric case and the mechanism with each agent in her own tier and a deterministic order in which agents are asked for evidence as the most asymmetric.

Tiers are useful to the principal if the agents are relatively asymmetric. Suppose, for example, that the agents all have the same distribution of values and all but one have the same cost, with the remaining agent having a much higher cost than the others. Suppose the principal uses a mechanism with only one tier and threshold v^* . Even if the agent with the high costs is last, we must make the threshold v^* very low and/or make her λ_i very high to give her an incentive to get evidence. Either is very costly to the principal. A low v^* makes it likely we stop before getting to this agent, giving the good to a low type of another agent. With λ_i very large, if none of the agents have virtual values above the threshold, the high-cost agent is likely to get the good even when there is an agent with a significantly higher value. With tiers, the principal can keep the threshold for the first $N - 1$ agents at v^* and only bring it down if none of these agents has a high enough value. This way, the principal gets the value of the lower threshold in incentivizing the high-cost agent but loses less on the first $N - 1$ agents than he would with a single threshold.

In principle, an optimal mechanism could differ from the description above in one more way. Specifically, it could be that the randomization over the next agent to ask for evidence depends nontrivially on the result of evidence received from previous agents. We call such a mechanism a *generalized tiered threshold mechanism*, or a generalized mechanism for short. We use the term *simple tiered threshold mechanism* or, more briefly, a simple mechanism for the class of mechanisms described above where the random order is not conditional in this way. We show that *every* optimal mechanism is a generalized mechanism. However, without loss of utility, the principal can restrict attention to simple mechanisms. More specifically, if a given mechanism is optimal, then there is a simple mechanism which is incentive compatible and yields the principal and every type of every agent the same expected payoff.

The key tools we use for these results are Weitzman's (1979) solution to the Pandora's box problem and Border's (1991) characterization of the set of feasible interim allocations. As we explain below, the Lagrangian for our problem takes exactly the form of Weitzman's problem if we treat the Lagrange multipliers as exogenous "preference parameters." This allows us to easily use Weitzman's characterization to show that every optimal mechanism is a generalized tiered threshold mechanism.

We then use Border’s theorem (appropriately reinterpreted) to show that we can restrict attention to simple mechanisms. As we explain in more detail later, the key observation is that the set of interim payoffs we can generate with simple mechanisms is the same as that for generalized mechanisms, so we do not need to consider the larger and more complex set.

In Section 2, we state the model. In Section 3, we characterize the optimal mechanism. In Section 4, we analyze the properties of the optimal mechanism, characterizing, in particular, the optimal random ordering of the agents. This section also provides some comparative statics results. We conclude with some extensions of the model in Section 5.

In the remainder of this section, we discuss the related literature. In addition to Weitzman (1979) and Border (1991), there are two related literatures. First, our work is connected to the literature on evidence, following the seminal work of Grossman (1981) and Milgrom (1981) as well as Green and Laffont’s (1986) analysis of mechanism design with evidence. See, for example, Glazer and Rubinstein (2004, 2006), Bull and Watson (2007), Deneckere and Severinov (2008), Hart, Kremer, and Perry (2017), and Ben Porath, Dekel, and Lipman (2019). In these papers, evidence is exogenous: the agent simply has certain evidence as a function of her type. By contrast, Ball and Kattwinkel (2025) and Ben Porath, Dekel, and Lipman (2026) do consider mechanism design when evidence can be acquired. These papers give some broad characterizations related to optimal one-agent mechanisms in these settings, but do not characterize optimal mechanisms for specific settings, as we do here.

Finally, there is a literature on mechanism design with information acquisition — see, for example, the survey of Bergemann and Välimäki (2006). In our model, the agent does not know her value until she acquires evidence, so evidence acquisition and information acquisition go hand-in-hand. The key difference between our work and these models, then, is exactly that the nature of the incentives to reveal the information acquired are different. With evidence, an agent is restricted in the misreports she can potentially get away with, so the honest-reporting constraints are different than in a model without evidence.

The most closely related papers are two papers in this literature, namely, Gershkov and Szentes (2009) and Crémer, Spiegel, and Zheng (2009). Both consider a principal and multiple agents. Gershkov–Szentes differs from our model in a few ways. First, they consider a *public* decision rather than a private allocation. The principal chooses a decision in $\{0, 1\}$ where all agents have the same state-dependent preferences between these options. In the optimal mechanism, the principal approaches agents in a random order to ask them to obtain information and provide it to him. Agents bear a private cost of obtaining information, as in our model. Because signals constitute evidence in

our model but not in theirs, they need to impose truth-telling constraints in addition to the obedience constraints in both models. They consider only the case where agents are symmetric and restrict attention to mechanisms that are ex post efficient. We do not need either of these restrictions.

Cr  mer, Spiegel, and Zhang, like us, consider a model where the principal is, in effect, allocating a single unit of a good. Unlike us, however, they allow monetary transfers and the principal’s objective function is revenue. Cr  mer, Spiegel, and Zhang assume that the principal controls the information acquisition and so can block an agent from getting information before the principal is ready for her to do so. Hence the principal can extract all ex ante surplus and so the principal’s problem becomes how to efficiently find a high-value buyer to sell the good to. Because of this, they can also use Weitzman’s (1979) results to characterize the optimal mechanism.³

2 Model

2.1 Setup

There are $N \geq 2$ agents and a principal. The principal has one unit of a good to allocate to an agent. The value to the principal of allocating the good to agent i is v_i . However, v_i is unknown to the principal or any agent at the outset. We assume that the common prior over v_i is given by cdf F_i with strictly positive density f_i over the support⁴ $[0, 1]$. We assume v_i ’s are independently distributed across agents. Aside from the assumption that the supports are the same, we impose no symmetry conditions on the distributions across agents. We sometimes let $V_i = [0, 1]$ denote the support of v_i and $V = [0, 1]^N = \prod_i V_i$.

Agent i can learn her value v_i at a cost $c_i \in (0, 1)$.⁵ If she pays this cost, she not only learns the realization of v_i but also obtains evidence enabling her to prove this realization to the principal. One can interpret “learning v_i ” as observing a verifiable signal which generates a certain conditional expectation of v_i . Agent i can then prove this conditional expectation by showing this signal to the principal. Since the principal is risk neutral, replacing v_i with this conditional expectation changes nothing.

We assume all agents want the good. Not including the cost of evidence acquisition,

³In a different use of Weitzman, see Khalfan and Vohra (2024) for an application to mechanisms with costly verification.

⁴We assume a bounded support for simplicity. As long as all relevant expectations exist, we could take the support to be $[0, \infty)$ with only minor technical complications.

⁵It is not difficult to extend our results to allow some agents to have costs above 1. See Section 5.

the agent's payoff is 1 if she receives the good, 0 otherwise.⁶ As we discuss in Section 5.2, the structure of the optimal mechanism is the same if we instead assume that the agent's payoff to receiving the good is some function $\varphi_i(v_i)$. This allows the agent's payoff to receiving the good to be correlated (positively or negatively) with the payoff to the principal of giving it to her. The agent's final payoff is the payoff from the allocation of the good minus c_i if she acquired evidence. The principal's payoff is independent of whether/which agents incur evidence costs and is equal to 0 if he keeps the good and v_i if he gives the good to agent i .

In some examples, it is natural to assume that agent i cannot “consume” the good without paying cost c_i . For example, consider departments in a university competing for a job slot. Suppose that the way departments prove their value to the university is by identifying their preferred candidate and showing his/her characteristics. It seems natural to suppose that even if a department were given the slot without needing to compete for it, they would have to pay the cost to identify whom to hire. We assume that paying the cost is not necessary for consumption in this sense, but our results and proofs would be unchanged if we assumed it is necessary. As we explain in Section 5.3, it is easier to extend our results to the case where the principal cares about the payoffs of the agents if we assume paying the cost is necessary for consumption.

The set of dynamic mechanisms available to the principal is quite complex. At each step, the principal can decide which agent or agents to ask for evidence, what (if anything) to tell them about what has happened so far, and how to react to the evidence they provide, if they do so. To keep the notation relatively simple, we restrict the class of mechanisms in a few ways that are clearly without loss of optimality for the principal.

First, we assume that the principal never asks more than one agent for evidence at a time. Because there is no discounting in our model, this is without loss for the principal as he can always ask one agent, then immediately afterward ask another.

Second, we assume that if an agent is asked to get evidence and refuses, then she is not given the good. By the Revelation Principle, we know that it is without loss of generality to consider mechanisms which induce agents to obey. Hence we may as well focus attention on mechanisms where agents are punished as severely as possible if they refuse to obey. In our model, the most severe possible punishment for an agent is not giving her the good.

Third, we assume no agent is asked for evidence more than once. The Revelation Principle says we can focus on mechanisms in which agents always obey. Since the

⁶Setting this payoff to 1 is a normalization. The substance of the assumption is that the expected value to agent i of receiving the good is strictly larger than the cost to her of obtaining evidence. If we let this payoff be φ_i and the cost be $\hat{c}_i$ instead, we can divide i 's payoff through by φ_i , making the payoff to getting the good equal to 1 and her cost equal to $\hat{c}_i/\varphi_i$.

principal cannot gain by getting the same evidence twice, he never asks any agent more than once.

Finally, we assume that the principal never provides any information to an agent upon asking her to obtain evidence. Put differently, each agent has (at most) one information set in the game the principal induces. To see that this is without loss for the principal, suppose instead that there are two different information sets for the agent. Then incentive compatibility requires that the agent's expected utility to obeying the principal is higher than her expected payoff to disobeying conditional on each information set. This implies but is not implied by the statement that the agent's expected utility to obeying is higher than her expected utility to disobeying unconditionally. Hence the principal weakly improves incentives by pooling these histories together.⁷

A dynamic mechanism specifies one of the following after every history. Either (a) the principal ends the process and keeps the good, (b) the principal ends the process and gives the good to an agent, or (c) the principal asks some agent to obtain and provide evidence.

Formally, a history is a sequence of agents and their responses to being asked for evidence. To be specific, a length n history is a sequence $((i_1, x_1), (i_2, x_2), \dots, (i_n, x_n))$ with the following properties. First, $i_n \in \{1, \dots, N\}$ for all n . That is, i_n is the agent who is the n th agent asked for evidence. Second, $x_n \in [0, 1] \cup \{R\}$. Here $x_n = R$ denotes the response of agent i_n to refuse to provide evidence. If $x_n \in [0, 1]$, then x_n is the value proved by agent i_n . Finally, if $i_k = i_\ell$, then $k = \ell$ — that is, no agent is asked for evidence twice.

Let H_n be the set of all length n histories and let $H = \cup_{n=0}^N H_n$ where we define $H_0 = \{e\}$, so e is the empty history. Note that there cannot be a history of length more than N .

A (pure) dynamic mechanism is a measurable function $d : H \rightarrow (\{0\} \times \{0, 1, \dots, N\}) \cup (\{1\} \times \{1, \dots, N\})$ satisfying the properties stated below. If $d(h) = (0, i)$, this means the principal ends the process on history h and gives the good to agent i (where $i = 0$ — that is, keeping the good — is possible). If $d(h) = (1, i)$, this means the principal continues the process on history h and asks agent i for evidence. Note that $d(h)$ cannot equal $(1, 0)$ — that is, the principal cannot ask himself for evidence.

We require d to satisfy the following properties. First, we require

$$d((i_1, x_1), (i_2, x_2), \dots, (i_n, x_n)) \neq (0, i_k)$$

if $x_k = R$ for any k . I.e., if i_k was asked for evidence and refused, she cannot get the

⁷Gershkov and Szentes (2009) use a similar argument. The earliest reference we know to this kind of reasoning is Myerson (1986).

good.

Second, we require that

$$d((i_1, x_1), (i_2, x_2), \dots, (i_n, x_n)) \neq (1, i_k)$$

for any $k \in \{1, \dots, n\}$. That is, the principal cannot ask any agent for evidence more than once.

Let $\hat{D}_p$ denote the set of pure dynamic mechanisms and let $\hat{D}$ denote the set of probability mixtures over $\hat{D}_p$ — that is, the set of random mechanisms.

2.2 The Principal's Problem

The following lemma shows that we can further restrict the set of dynamic mechanisms to those that satisfy a property we call *no free lunch*. We say a mechanism satisfies no free lunch if it never gives the good to an agent who has not provided evidence. It is easy to see that the principal may as well choose a mechanism satisfying this property as he does not pay the costs — if he were about to give the good to an agent without evidence, he can always ask for evidence, promising her the good iff she shows *any* evidence to him. Since $c_i < 1$, the agent will obey and the principal's payoff is unchanged. The following lemma establishes the more substantial point that *every* optimal dynamic mechanism satisfies this property.⁸

Lemma 1. *Fix any incentive compatible dynamic mechanism violating no free lunch on a set of histories with strictly positive probability. Then there is another incentive compatible mechanism which gives the principal a strictly higher expected payoff. So every optimal incentive compatible mechanism satisfies no free lunch up to sets of measure zero.*

Proof. Suppose H^* is a positive probability set of histories on which the principal gives the good to agent i with positive probability even though i has not provided evidence. There are two cases. First, suppose every history in H^* has the property that the principal previously received evidence from at least one agent. (If H^* does not have this property but some positive measure subset does, we can replace H^* with this subset.) For each history $h \in H^*$, let $\bar{v}(h)$ denote the highest value which some other agent has previously proven to the principal. Because this set of histories has positive measure, there is a

⁸The strict gain established in this lemma makes nontrivial use of the continuum of types. With finite type sets, there could be a strictly positive probability that the principal has checked all but one agent and all those checked have the lowest possible value. In this situation, the principal cannot strictly gain by asking the last agent for evidence, so switching from a mechanism that violates no free lunch to one satisfying it leaves the principal indifferent.

natural number n such that the set of histories $h \in H^*$ with $\bar{v}(h) > 1/n$ has strictly positive measure.

Let $\bar{v}_i$ be defined by $1 - F_i(\bar{v}_i) = c_i$. Because $c_i < 1$ for all i and because F_i is continuous for all i , we know that $\bar{v}_i > 0$. Fix some $\varepsilon \in (0, \bar{v}_i)$.

Change the mechanism only on histories in H^* as follows. With the probability the original mechanism gave the good to agent i , the principal instead asks i for evidence. If i does not provide evidence, the principal keeps the good. If i gives evidence showing either $v_i \geq \bar{v}_i - \varepsilon$ or $v_i \geq \bar{v}(h)$, the principal gives the good to i . If i 's evidence shows that $v_i < \bar{v}_i - \varepsilon$ and $v_i < \bar{v}(h)$, the principal gives the good to the agent who proved value $\bar{v}(h)$.

It is easy to see that i 's probability of receiving the good if she provides evidence is strictly larger than her cost, so it is strictly optimal for her to provide evidence. Also, every other agent's incentive to provide evidence is at least as large as before since any agent who obtains evidence now gets rewarded with the good at least as often. Hence the new mechanism is incentive compatible. Clearly, the probability that the principal gains by giving the good to an agent with a value higher than v_i is strictly positive, so the principal's expected payoff to the alternative mechanism is strictly larger. Hence the original mechanism was not optimal.

Second, suppose that the set of histories for which the principal gives the good to i with positive probability without obtaining evidence from i has positive probability, but the subset of these histories on which some other agent has previously provided evidence has probability zero. By incentive compatibility, this means that on these histories, no other agent has been asked for evidence. In this case, fix any agent $j \neq i$ and replace $\bar{v}(h)$ in the argument above with $E(v_j)$. The change in the mechanism now has no effect on the incentives of other agents to obtain evidence since the change only takes place when none of them have been asked to do so. Hence the new mechanism is again incentive compatible and improves the principal's payoff, implying that the original mechanism was not optimal. ■

Let D_p denote the set of pure dynamic mechanisms $d \in \hat{D}_p$ satisfying no free lunch — i.e., those such that $d(h) = (0, i)$ only if $h = ((i_1, x_1), (i_2, x_2), \dots, (i_n, x_n))$ where $i = i_k$ for some k . Let D denote the set of probability mixtures over D_p .

By the Revelation Principle, it is without loss of generality to focus attention on mechanisms in which agents find it optimal to obey the principal. That is, once we restrict to incentive compatible mechanisms, we know that the relevant histories will be ones where agents who are asked for evidence do provide it. Given such a dynamic mechanism, we can compute the outcome under the mechanism as a function of the

profile of types v .⁹

Let $P(d) = (P_1(d), \dots, P_N(d))$ denote the allocation probabilities induced by dynamic mechanism d . That is, for each d , $P_i(d)$ is a measurable function mapping V to $[0, 1]$ where $\sum_i P_i(v | d) \leq 1$ for all $v \in V$. Let $e_i(d)$ denote the probability agent i is asked for evidence in mechanism d .

We can now state the principal's maximization problem. The principal's objective function is $E_v [\sum_i P_i(v | d)v_i]$.

The constraints are the N incentive compatibility constraints. One might expect these constraints to be very complex since they say that *conditional on being asked for evidence*, an agent finds it optimal to obey. This conditioning depends in a complex way on the dynamic mechanism since we have to identify the set of histories on which this agent might be asked for evidence. Fortunately, we are able to bypass this complexity by expressing the incentive compatibility constraint at the ex ante stage. Recall that each agent has (at most) one information set in the mechanism. Hence we can write the incentive compatibility constraint as requiring that the ex ante optimal strategy for the agent is to obtain evidence and provide it if asked for it. Note that if the agent obeys the principal, then her expected payoff in mechanism d must be $E_v P_i(v | d) - c_i e_i(d)$.

To pin down the agent's deviation payoff, first, consider the deviation strategy where the agent does get evidence when asked but does not always report it. In this case, her expected costs of evidence acquisition are still $c_i e_i(d)$, but she must receive the good weakly less often. Hence obeying the principal must give a weakly higher payoff than this. Second, consider the deviation strategy of not getting evidence. In this case, she gets a payoff of 0 whether she is asked for evidence (because she disobeys) or not (because of no-free-lunch) and hence her expected payoff is 0. Therefore, we can write the incentive compatibility constraint for agent i as $E_v P_i(v | d) - c_i e_i(d) \geq 0$.

Hence we can state the principal's optimization problem as follows. We say that $d \in D$ is *optimal* if it solves the problem

$$\max_{d \in D} E_v \left[\sum_i P_i(v | d)v_i \right]$$

subject to

$$E_v [P_i(v | d)] - c_i e_i(d) \geq 0, \quad \forall i.$$

Let D^* denote the set of optimal d 's. In the next section, we characterize this set.

⁹We omit the precise definition as it will not be needed. Briefly, one can iteratively define the probability distribution over realized histories given d as a function of the profile v . This then determines the probability distribution over outcomes.

3 Characterizing the Optimal Mechanism

We characterize the optimal mechanism in two steps. The first uses Weitzman (1979), the second Border (1991). First, we briefly summarize Weitzman's results.

3.1 Weitzman

Weitzman considers the following problem, simplified here to more easily line up with our problem. There is a searcher who faces N "boxes." Box i has a certain monetary prize $x_i \geq 0$ in it where x_i is distributed according to a distribution $\hat{F}_i$. Prizes are independently distributed across boxes. There is a cost, $\hat{c}_i$, to opening box i . At each point in the search process, the searcher decides between quitting and taking no box, quitting and taking some box she has previously opened, or opening a box she has not yet opened. The searcher's payoff is the prize in the box she takes (or 0 if she takes no box) minus the accumulated costs of the boxes she has opened.

Weitzman characterizes the set of optimal search procedures as follows. For each box i , define an index, r_i , by

$$\hat{c}_i = E_{x_i} \max\{x_i - r_i, 0\}.$$

Intuitively, r_i is that value such that the searcher would be indifferent between stopping with value r_i or opening box i and then quitting with the larger of r_i and the prize in box i . For simplicity, for this discussion, we assume $r_i > 0$ for all i , as the analog of this property will necessarily hold for our use of Weitzman's result. Given this, a search procedure is optimal iff it takes the following form. First, the searcher opens any box i_1 with the highest index — i.e., such that $r_{i_1} = \max_j r_j$. If there is more than one such box, *any* randomization is optimal. If the prize in box i_1 , x_{i_1} , satisfies $x_{i_1} > \max_{j \neq i_1} r_j$, then the searcher stops and takes box i_1 . In our problem, the analog of the x_i 's will be continuously distributed, so we do not need to consider what happens if $x_{i_1} = \max_{j \neq i_1} r_j$. If $x_{i_1} < \max_{j \neq i_1} r_j$, the searcher opens any box i_2 with the highest index of the remaining boxes — i.e., such that $r_{i_2} = \max_{j \neq i_1} r_j$. The searcher continues in this fashion, comparing the largest prize found so far in any box to the maximum index of the unopened boxes. If the largest prize is strictly above the highest index among the unopened boxes, the searcher stops and takes the corresponding box. Otherwise, she continues and opens any unopened box with the largest possible index. If she opens all boxes, she takes the one with the largest prize.

The following lemma will link Weitzman's result to our problem. Recall that D^* is the

set of optimal d 's. Given $\lambda \in \mathbf{R}_+^N$, let $D^{**}(\lambda)$ be the set of maximizers of the Lagrangian

$$\mathcal{L} = \mathbb{E}_v \left[\sum_i P_i(v \mid d) v_i \right] + \sum_i \lambda_i [\mathbb{E}_v P_i(v \mid d) - c_i e_i(d)]$$

and let D^{**} denote the set of incentive compatible $d \in D^{**}(\lambda)$ for some λ such that $\lambda_i [\mathbb{E}_v P_i(v \mid d) - c_i e_i(d)] = 0$ for all i , the complementary slackness conditions.

Lemma 2. $D^* = D^{**}$.

In other words, strong duality holds. This follows from the fact that the set of (P, e) that can be generated by a mechanism is convex and that payoffs are linear in (P, e) . The proof of this result is relatively standard but is contained in the supplemental appendix for the convenience of the reader.

We can rewrite the Lagrangian as

$$\mathbb{E}_v \left[\sum_i P_i(v \mid d) (v_i + \lambda_i) - \lambda_i c_i e_i(d) \right].$$

Let d^* denote an optimal mechanism. By Lemma 2, there exist Lagrange multipliers λ^* such that d^* solves the problem

$$\max_{d \in D} \mathbb{E}_v \left[\sum_i P_i(v \mid d) (v_i + \lambda_i^*) - \lambda_i^* c_i e_i(d) \right].$$

This is almost exactly Weitzman's problem. Think of agent i as box i where the prize in box i is $v_i + \lambda_i^*$ and the cost of opening box i is $\lambda_i^* c_i$. Think of $P_i(v \mid d)$ as the probability of choosing box i under search procedure d when the prizes are given by $(v_1 + \lambda_1^*, \dots, v_N + \lambda_N^*)$ and $e_i(d)$ as the probability of opening box i under search procedure d .

The only difference between this problem and Weitzman's is that a mechanism in our problem must specify what to do on a history where some agent has been asked for evidence and refused to comply. In Weitzman's problem, a search procedure is not defined on such a history as boxes cannot refuse to be opened. Let H^* denote the set of all possible histories in our problem with the property that no agent has ever refused when asked for evidence. Then the set of search procedures in Weitzman is exactly our set of dynamic mechanisms when we restrict the set of histories to H^* . Because we can restrict attention to incentive compatible mechanisms, no agent will ever refuse to provide evidence if asked. Hence the histories we exclude when considering H^* are payoff irrelevant.

Hence d^* must be in the class of search procedures Weitzman identifies. The index for box/agent i , which we denote by $\hat{v}_i + \lambda_i^*$, is defined by

$$\lambda_i^* c_i = E_{v_i + \lambda_i^*} \max\{v_i + \lambda_i^* - (\hat{v}_i + \lambda_i^*), 0\} = E_{v_i} \max\{v_i - \hat{v}_i, 0\}.$$

Partition the set of agents into tiers, $\mathcal{I}_1, \dots, \mathcal{I}_K$ where the agents in $\mathcal{I}_1$ have the largest index, those in $\mathcal{I}_2$ have the next largest index, etc. For tier $\mathcal{I}_k$, let v_k^* denote the common index for the agents in that tier.

Then the optimal procedure starts with the agents in tier 1, asking them for evidence in some order. This order is arbitrary in Weitzman and so could depend on the specific agents previously asked for evidence and/or the evidence they show. We refer to the value of the prize in the box, $v_i + \lambda_i$, as i 's *virtual value*. If i 's virtual value is larger than v_1^* , then we stop and take that box/give the good to that agent. Otherwise, we continue to some other box/agent in tier 1. After checking the last agent in tier 1, the relevant comparison is to the common index for the second tier, v_2^* . Thus if the agent with the highest virtual value in tier 1 is above v_2^* , this agent gets the good and otherwise we continue to tier 2, etc. If we end up opening all the boxes/getting evidence from all the agents, we take the highest prize — that is, we give the good to the agent with the highest virtual value.

We call a mechanism of this form a *generalized tiered threshold mechanism* or a generalized mechanism for brevity. More specifically, a generalized mechanism divides the agents into tiers, using a random, possibly history-dependent order for asking them for evidence. We stop and give the good to agent i after she is checked iff $v_i + \lambda_i$ is larger than a tier-specific constant or if we have checked all agents in i 's tier, $v_i + \lambda_i$ exceeds the threshold for the next tier, and is the largest virtual value seen so far. Finally, if all agents are checked, the one with the highest virtual value receives the good.

Hence the discussion above proves the following result:

Theorem 1. *Every optimal dynamic mechanism is a generalized tiered threshold mechanism, up to sets of measure zero.*

Note that Weitzman's theorem implies that the optimal mechanism has the form of a generalized tiered threshold mechanism. However, his results do not identify the randomization over the order of checking, not even whether it varies with previous observations by the principal. In Weitzman, if two boxes have the same index, then any randomization over which to check first is equally good, including randomizations that depend on the prizes in previously opened boxes. For our model, though, these randomizations are crucial for incentive compatibility. As discussed in the introduction, all else equal, an agent who is asked earlier for evidence is less likely to receive the good conditional on being asked. Hence an agent who is more likely to be asked early has

less incentive to obey when asked for evidence. In short, Weitzman's theorem identifies the form of the optimal mechanism for some profile of λ_i 's without identifying anything about the randomizations involved. The next step is to identify the λ_i 's and characterize the randomizations.

The space of possible randomizations is very large. In principle, any randomization over which agent to check next as a function of the history so far could be part of an optimal mechanism.

3.2 Interim Equivalence and Border

Our next step is to use Border (1991) to both restrict the set of possible randomizations and to focus on objects simpler than randomizations. The key to this is showing that for any generalized tiered threshold mechanism, we can find a simpler mechanism which is equivalent in the following sense.

Definition 1. *Mechanisms d and d' are interim-equivalent if for all i and almost all $v_i \in [0, 1]$, $E_{v_{-i}} P_i(v_i, v_{-i} \mid d) = E_{v_{-i}} P_i(v_i, v_{-i} \mid d')$ and if for all i , $e_i(d) = e_i(d')$.*

Because the agents are not endowed with private information, it would be natural to call two mechanisms equivalent if they gave the principal and every agent the same ex ante expected payoff. We use the stronger notion of interim-equivalence both because it gives a stronger result and because the interim comparison is convenient for proving our results.

For brevity, we define the usual interim (or reduced form) probabilities by $p_i(v_i \mid d) = E_{v_{-i}} P_i(v_i, v_{-i} \mid d)$.

The following lemma shows the significance of interim-equivalence.

Lemma 3. *If d is an optimal incentive-compatible mechanism and d' is interim-equivalent to it, then d' is also an optimal incentive-compatible mechanism and every type of every agent obtains the same payoff in d' as in d .*

Proof. Since d is incentive compatible, we have

$$0 \leq E_v P_i(v \mid d) - c_i e_i(d) = E_{v_i} p_i(v_i \mid d) - c_i e_i(d)$$

for all i . Hence, since d' is interim-equivalent to d , we have

$$E_{v_i} p_i(v_i \mid d') - c_i e_i(d') \geq 0$$

for all i , so d' is also incentive compatible. More generally, the payoff to agent i of type v_i in mechanism d is

$$\mathbb{E}_{v_{-i}} P_i(v_i, v_{-i} \mid d) - c_i e_i(d) = p_i(v_i \mid d) - c_i e_i(d) = p_i(v_i \mid d') - c_i e_i(d'),$$

so every type of every agent is indifferent between the two mechanisms.

Finally, we can write the payoff of the principal under d as

$$\mathbb{E}_v \left[\sum_i P_i(v \mid d) v_i \right] = \sum_i \mathbb{E}_{v_i} [\mathbb{E}_{v_{-i}} P_i(v_i, v_{-i} \mid d) v_i] = \sum_i \mathbb{E}_{v_i} [p_i(v_i \mid d) v_i].$$

Since d' is interim-equivalent to d , the principal's payoff is the same under d and d' , so if d is optimal, d' must be optimal as well. ■

Lemma 3 says we can identify optimal mechanisms (at most) up to interim-equivalence. Hence we may as well focus on a convenient selection from the interim-equivalent optimal mechanisms. The next lemma gives an easy way to identify certain pairs of interim-equivalent mechanisms.

Lemma 4. *Fix mechanisms d and $\hat{d}$ in $D^{**}(\lambda)$ satisfying $e_i(d) = e_i(\hat{d})$ for all i . Then d and $\hat{d}$ are interim-equivalent.*

Proof. We show that we can write $p_i(v_i \mid d)$ entirely as a function of the λ 's and e_i . Given this, if two mechanisms have the same λ 's and e 's, they must be interim-equivalent.

So fix any $\lambda \in \mathbf{R}_+^N$ and any $d \in D^{**}(\lambda)$. The proof of Theorem 1 shows that this must be a generalized tiered threshold mechanism. Fix the e 's generated by this mechanism.

The proof of Theorem 1 shows that we can define the index for i entirely from λ_i . Specifically, it is $\hat{v}_i + \lambda_i$ where $\hat{v}_i$ is defined by

$$\lambda_i c_i = \int_{\hat{v}_i}^1 (v - \hat{v}_i) f_i(v) dv.$$

It is easy to see that $\hat{v}_i$ is uniquely determined by λ_i .

Given the profile $(\hat{v}_1 + \lambda_1, \dots, \hat{v}_N + \lambda_N)$ and $e = (e_1, \dots, e_N)$, we can compute $p_i(v_i \mid d)$ for any i and any v_i as follows. First, if $v_i \geq \hat{v}_i$, then agent i receives the good if and only if she is asked for evidence, so $p_i(v_i \mid d) = e_i(d)$. Second, if $v_i < \hat{v}_i$, we have

$$p_i(v_i \mid d) = \prod_{j \neq i \mid \hat{v}_j + \lambda_j \geq v_i + \lambda_i} F_j(v_i + \lambda_i - \lambda_j).$$

To see this, first, consider $v_i \geq \hat{v}_i$. By no-free-lunch, i does not receive the good if she is not asked for evidence. If she is asked for evidence and $v_i \geq \hat{v}_i$, then $v_i + \lambda_i \geq \hat{v}_i + \lambda_i$, so i 's virtual value is higher than her Weitzman index. Since she is being asked for evidence, the Weitzman index for every agent who has not yet been asked for evidence must be below $\hat{v}_i + \lambda_i$, so her virtual value is above the relevant threshold. Hence she receives the good. In short, if $v_i \geq \hat{v}_i$, we have $p_i(v_i \mid d) = e_i(d)$.

So suppose $v_i < \hat{v}_i$, so i 's virtual value is below her index. In this case, i receives the good only if all the other agents who are checked have virtual values below hers. More precisely, note that any agent j with $\hat{v}_j + \lambda_j > v_i + \lambda_i$ will be checked before agent i is given the good. So for i to receive the good when her value is v_i , it must be the case that all such j have $v_j + \lambda_j < v_i + \lambda_i$. The expression above gives the probability of this event.

Hence $p_i(v_i \mid d)$ is uniquely identified by the λ 's and e 's generated by d . ■

We now use this result to show that for every generalized mechanism, there is what we call a *simple tiered threshold mechanism* (or *simple mechanism* for short) which is interim-equivalent to it. A simple mechanism differs from a generalized mechanism only in that the randomization over the order in which agents in a given tier are asked is independent of the history to that point. To be more specific, for each tier $\mathcal{I}_k$, there is a probability distribution, denoted O_k over the set of linear orders over $\mathcal{I}_k$ where we interpret a typical such order, $\succ_k$, by saying that if $i \prec_k j$, then i goes before j . In other words, for each tier, we have a fixed, history-independent, one-time randomization determining the order in which agents in that tier are asked for evidence.

We use Border's theorem for this purpose. The idea is to show that given any λ 's and any e 's generated by some generalized mechanism, there exists a simple mechanism with the same λ 's generating the same e 's. By Lemma 4, the claim then follows.

To do this, we characterize the set of e 's that can be generated by an optimal mechanism given the λ 's. To see the issue, suppose tier 1 consists of agents 1 and 2, that $\hat{v}_1 = \hat{v}_2 \in (0, 1)$, and that $e_1 = e_2 = 1$. From the above, if $v_1 \geq \hat{v}_1$ and $v_2 \geq \hat{v}_2$, both agents 1 and 2 have virtual values above the threshold for tier 1. Hence both get the good iff they are asked for evidence. But we are hypothesizing that both are asked for evidence with probability 1. So this implies both get the good when $v_1 \geq \hat{v}_1$ and $v_2 \geq \hat{v}_2$ which is impossible.

In other words, given the λ 's, not every $(e_1, \dots, e_N)$ can be generated by some choice of randomization over the order of asking agents for evidence. The issue may seem different, but it turns out to be related to Border's (1991) characterization of the set of interim allocation functions which are feasible in the sense that they can be generated by some allocation functions. Border's result covered symmetric distributions and subsequent

work extended his results in many directions — see, for example, Mierendorff (2011) or Che, Kim, and Mierendorff (2013). Here we give some clearly necessary conditions on the e_i 's in the spirit of the Border conditions. In the Appendix, we show that these conditions are sufficient.

In the example above, the e 's are not feasible because the implied interim allocation probabilities violate the usual Border conditions. However, the conditions we require go beyond the usual Border conditions applied to the interim allocation probabilities. To see this, take the same example as above but suppose $e_1 = e_2 = 0$. In this case, the formula given above for computing the interim allocation given the λ 's and e 's would say that the allocation gives the good to the agent with the higher value if this value is less than $\hat{v}$ and to neither agent otherwise. This interim allocation is certainly feasible in the usual model and so will satisfy the usual Border conditions. However, it is not feasible in our model as the principal cannot allocate the good to the agent with the higher value if he never obtains evidence. Thus the restrictions we require are not just that the induced interim probabilities satisfy the usual Border conditions.

Lemma 5. *Fix $\lambda = (\lambda_1, \dots, \lambda_N)$ and the associated $\hat{v}_1, \dots, \hat{v}_N$. If there exists a generalized tiered threshold mechanism $d \in D^{**}(\lambda)$ which generates $(e_1(d), \dots, e_N(d)) = (e_1, \dots, e_N)$, then the following conditions hold. First,*

$$\sum_{i \in \mathcal{I}_k} e_i [1 - F_i(\hat{v}_i)] = \left[\prod_{i \in \mathcal{I}^{k-1}} F_i(v_k^* - \lambda_i) \right] \left[1 - \prod_{i \in \mathcal{I}_k} F_i(\hat{v}_i) \right], \quad (1)$$

where the first term on the right-hand side is defined to be 1 for $k = 1$ and otherwise

$$\mathcal{I}^{k-1} = \bigcup_{\ell=1}^{k-1} \mathcal{I}_\ell.$$

Second, for all $\mathcal{I} \subset \mathcal{I}_k$,

$$\sum_{i \in \mathcal{I}} e_i [1 - F_i(\hat{v}_i)] \leq \left[\prod_{i \in \mathcal{I}^{k-1}} F_i(v_k^* - \lambda_i) \right] \left[1 - \prod_{i \in \mathcal{I}} F_i(\hat{v}_i) \right], \quad (2)$$

If the mechanism satisfies the complementary slackness conditions $\lambda_i [E_v P_i(v \mid d) - c_i e_i(d)] = 0$ for all i , then $e_i = 1$ for any i with $\lambda_i = 0$.

Furthermore, for any e satisfying (1), (2), and the property that $e_i = 1$ for any i with $\lambda_i = 0$, there is a simple tiered threshold mechanism $d' \in D^{**}(\lambda)$ that generates $e_i(d') = e_i$ for all i .

Because of their similarity to the conditions identified by Border (1991), we refer to equations (1) and (2) and the restriction that $e_i = 1$ if $\lambda_i = 0$ as the Border conditions.

Proof. The proof that any e generated by a generalized tiered threshold mechanism must satisfy equations (1) and (2) is straightforward. First, consider tier 1. For $k = 1$, the first Border condition, equation (1), says

$$\sum_{i \in \mathcal{I}_1} e_i [1 - F_i(\hat{v}_i)] = 1 - \prod_{i \in \mathcal{I}_1} F_i(\hat{v}_i).$$

To understand the left-hand side, consider an agent i in tier 1 with $v_i \geq \hat{v}_i$ (equivalently, $v_i + \lambda_i \geq \hat{v}_i + \lambda_i$). By the definition of a generalized tiered threshold mechanism, such an agent gets the good if and only if she is asked for evidence. Hence the left-hand side is the probability that the good goes to some agent i in tier 1 with a value in this range. However, the definition of a generalized tiered threshold mechanism also says that if there is some agent i in tier 1 with a value in this range, then the good *must* go to such an agent. Since the right-hand side is the probability that the realized v has at least one tier 1 agent with a value in this range, we see that the left-hand side and right-hand side must be equal.

Continuing with $k = 1$, the second Border condition, equation (2), says

$$\sum_{i \in \mathcal{I}} e_i [1 - F_i(\hat{v}_i)] \leq 1 - \prod_{i \in \mathcal{I}} F_i(\hat{v}_i),$$

for all $\mathcal{I} \subseteq \mathcal{I}_1$. To see that this must hold, note that the left-hand side is the probability that an agent $i \in \mathcal{I}$ has $v_i \geq \hat{v}_i$, is asked for evidence, and hence receives the good, while the right-hand side is the probability that some agent $i \in \mathcal{I}$ has $v_i \geq \hat{v}_i$. Hence the left-hand side must be smaller than the right.

Similarly, consider $k = 2$, where the first Border condition says that

$$\sum_{i \in \mathcal{I}_2} e_i [1 - F_i(\hat{v}_i)] = \left[\prod_{i \in \mathcal{I}_1} F_i(v_2^* - \lambda_i) \right] \left[1 - \prod_{i \in \mathcal{I}_2} F_i(\hat{v}_i) \right].$$

Now the left-hand side is the probability that an agent i in tier 2 has a value $v_i \geq \hat{v}_i$, is asked for evidence, and hence receives the good. We know that this happens if and only if all the tier 1 agents have virtual values below the tier 2 threshold v_2^* and some tier 2 agent has $v_i \geq \hat{v}_i$. A tier 1 agent i has virtual value below the tier 2 threshold iff $v_i + \lambda_i < v_2^*$ or $v_i < v_2^* - \lambda_i$. So the right-hand side is exactly the probability a tier 2 agent with virtual value above the tier 2 threshold gets the good.

For $k = 2$, the second Border condition says that

$$\sum_{i \in \mathcal{I}} e_i [1 - F_i(\hat{v}_i)] \leq \left[\prod_{i \in \mathcal{I}_1} F_i(v_2^* - \lambda_i) \right] \left[1 - \prod_{i \in \mathcal{I}} F_i(\hat{v}_i) \right],$$

for every $\mathcal{I} \subseteq \mathcal{I}_2$. In this case, the left-hand side is the probability that some agent $i \in \mathcal{I}$ has a value above $\hat{v}_i$ and gets the good, while the right-hand side is the obviously necessary condition that no agent in tier 1 gets the good before we get to tier 2 and that there is some agent $i \in \mathcal{I}$ with $v_i \geq \hat{v}_i$. Hence this inequality is necessary.

The proof for other tiers is analogous. The implications of complementary slackness and that any e satisfying the Border conditions can be generated by a simple tiered threshold mechanism are shown in the Appendix. ■

Corollary 1. *For any optimal generalized tiered threshold mechanism d , there is a simple tiered threshold mechanism d' which is interim-equivalent to it. Hence there is always an optimal mechanism which is a simple mechanism.*

To see why the corollary holds, note that Lemma 5 implies that the e generated by any $d \in D^{**}(\lambda)$ can also be generated by a simple tiered threshold mechanism $\hat{d} \in D^{**}(\lambda)$. By Lemma 4, then, for any generalized tiered threshold mechanism, there is an interim-equivalent simple mechanism.

Summarizing, we have shown

Theorem 2. *An incentive-compatible mechanism d is optimal if and only if it is a generalized tiered threshold mechanism with tiers defined by $\hat{v}_1 + \lambda_1, \dots, \hat{v}_N + \lambda_N$ where $\lambda_i \geq 0$ for all i ,*

$$\lambda_i c_i = \int_{\hat{v}_i}^1 (v - \hat{v}_i) f_i(v) dv, \quad \forall i,$$

$$\lambda_i [E_{v_i} p_i(v_i \mid \lambda, e_i) - c_i e_i] = 0, \quad \forall i,$$

and the Border conditions hold, where $e_i = e_i(d)$ for all i .

Furthermore, given any optimal mechanism d , there is a simple tiered threshold mechanism which is also optimal and yields every type of every agent the same expected payoff.

3.3 Tiers

A tempting but incorrect intuition suggests that the complexities of identifying these randomizations can “typically” be avoided. We only need to identify the random order in which agents are asked when agents are in the same tier. Agents are only in the same tier when they have the same index. It is tempting to suspect that for “generic” c_i ’s and F_i ’s, indices never tie, so, in this sense, randomization is “almost always” irrelevant.

This is not correct. Recall that the index, $\hat{v}_i + \lambda_i$, is defined by

$$\lambda_i c_i = \int_{\hat{v}_i}^1 (v - \hat{v}_i) f_i(v) dv.$$

Hence the index depends on the exogenous c_i and F_i , but also on the endogenous λ_i . This endogeneity leads to ties in the indices and hence tiers with more than one agent. In fact, “similar enough” agents *must* be in the same tier, so that ties are not “measure zero.”

To see the intuition, first consider the (nongeneric) symmetric agent case. Suppose there are two agents with the same F_i ’s and the same c_i ’s and assume the incentive compatibility constraint is binding¹⁰ for both agents. In this case, the optimal mechanism *must* randomize 50–50 over who to start with.

To see this, suppose we instead put these agents in different tiers, asking 1 for evidence first and then potentially asking 2. Intuitively, the optimal mechanism subject to this ordering must have both incentive constraints binding. For this reason, each must be indifferent between getting evidence and not doing so conditional on being asked. Because 1 has the disadvantage of going first, this means that 1 must have an advantage over 2 in the sense that 1 gets the good when both are checked whenever $v_1 + \Delta > v_2$ for some $\Delta > 0$. Clearly, if we started with 2 first instead, 2 would be favored in this sense by the same Δ . Similarly, if we randomize 50–50 over these two mechanisms *informing the agents of the outcome of the randomization*, then whoever is selected to go first has the same Δ advantage.

But suppose we randomize 50–50 and *don’t* tell the agents who was selected to go first. In this case, we can reduce Δ and maintain both incentive constraints. Intuitively, a lower Δ hurts the agent when she goes first, but helps her when she’s second. With a 50–50 randomization, one effect exactly offsets the other. The smaller Δ means that the principal is able to take the higher value more often and hence receives a strictly higher payoff. To minimize the loss associated with favoring one agent over the other, the strictly best option for the principal is to randomize 50–50 over which agent is first.

Given that a 50–50 randomization is the unique solution for identical agents, it should not be surprising that nearby randomizations are the unique solution for nearly identical agents. Hence ties are not “nongeneric.”

¹⁰Here and elsewhere, we say a constraint is binding if the principal would obtain a strictly higher payoff in its absence. In nongeneric situations, we can have the constraints holding with equality at the solution and yet not binding in this sense. For example, with two symmetric agents with cost equal to $1/2$, we obtain the first–best even though both agents receive zero utility and hence have incentive constraints holding with equality. It is not hard to show that if the incentive constraint for i is binding in our sense, then her payoff is zero at the optimum, even though the converse can be violated. Also, one can show that the incentive constraint for agent i is binding in our sense if and only if $\lambda_i > 0$ at the optimum.

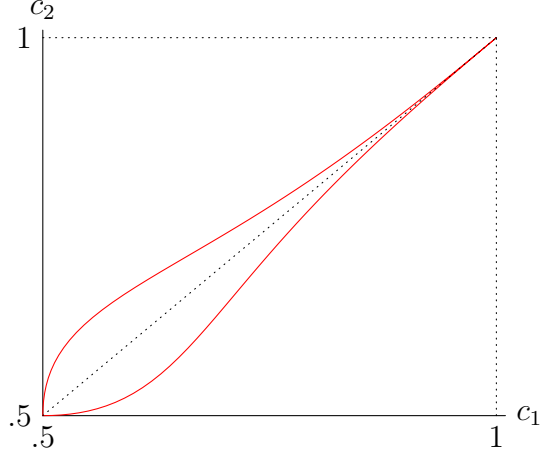


Figure 1: (c_1, c_2) values with one tier when $N = 2$, $v_1, v_2 \sim U[0, 1]$.

The figure above illustrates. Assuming two agents, each with $v_i \sim U[0, 1]$, the area between the red curves is the set of (c_1, c_2) in the range $[\frac{1}{2}, 1]^2$ where the two agents have the same index in the optimal mechanism. As the intuition above suggests, it is a non-negligible set of parameters around symmetry.

4 Properties of the Optimal Mechanism

In this section, we describe properties of optimal mechanisms. Section 4.1 characterizes the optimal (random) ordering of agents. Section 4.2 discusses the symmetric case where $c_i = c_j$ and $F_i = F_j$ for all i, j , in more detail, particularly focusing on comparative statics. Finally, Section 4.3 comments briefly on comparative statics in the general case.

4.1 Optimal Ordering

In the symmetric case, the agents are identical, so there is no reason for the principal to prefer one order for seeking evidence over another. When agents are asymmetric, what determines the optimal randomization over the order?

Intuitively, agents who are later in the order are more protected from competition. These agents are only asked for evidence when the values of the earlier agents are relatively low, so being asked is a good sign for them about the competition they face. This suggests the principal will put “stronger” agents earlier. Theorem 3 shows that this intuition is correct.

We say that agent i is *stronger than* agent j if $c_i \leq c_j$ and F_i (weakly) first-order stochastically dominates F_j . Note that one or both of these comparisons can be an equality relation — i.e., agents with the same cost and same distribution are each stronger than the other.

Theorem 3. *If i is stronger than j , then in an optimal mechanism, we have $e_i \geq e_j$.*

The following corollaries elucidate the implications of this result.

Corollary 2. *If i is stronger than j and $\max\{e_i, e_j\} > 0$, then $\hat{v}_i + \lambda_i \geq \hat{v}_j + \lambda_j$, so if they are in different tiers, i is in a higher tier than j . Hence, if the optimal order is deterministic, i is asked before j .*

If i and j have different indices and so are in different tiers, then $e_i \geq e_j$ implies that i must be in the higher tier. This is because agents in a tier are not asked for evidence until *all* agents in *all* higher tiers have been asked. The indices are not unique for agents with $e_i = 0$, requiring the focus on positive evidence probabilities.

Corollary 3. *If i is stronger than j and j is stronger than i , then $e_i = e_j$. In this case, the optimal order cannot be deterministic if their incentive constraints are binding.*

If i and j have the same costs and same distribution, each is stronger than the other, so Theorem 3 implies that $e_i = e_j$. If the order is deterministic, this means that i is asked for evidence if and only if j is also asked. This cannot be optimal if the incentive constraints for i and j are both binding, so the optimal order *cannot* be deterministic. Note that if both incentive constraints are slack, both agents are asked for evidence with probability 1 and the principal does not care whether he asks both at the same time or in some order.

Corollary 4. *If i is stronger than j and i is excluded in the sense that she never receives the good, then j is excluded.*

To see this, note that i is excluded if and only if $e_i = 0$. If $e_i = 0$, then no-free-lunch implies that i never receives the good and hence is excluded. Conversely, if i never receives the good, then her incentive constraint becomes $-c_i e_i \geq 0$, implying $e_i = 0$. If i is stronger than j and i is excluded, then we have $e_i = 0$ and hence $e_j = 0$.

4.2 Symmetric Mechanisms

When the agents are symmetric, the analysis simplifies greatly. In this case, the principal doesn't care *which* agent he asks for evidence and only needs to decide *whether* to seek evidence at any given history.

More formally, in the symmetric setting, there must be an optimal mechanism which is symmetric in the sense that it treats the agents identically. Thus λ_i , $\hat{v}_i$, and e_i are all independent of i . In this subsection, we drop the i subscripts on these variables.

Since all agents have the same $\hat{v} + \lambda$, clearly, there is only one tier and one threshold. In this case, there is no point distinguishing between virtual values, $v_i + \lambda$, and actual values, v_i , since the λ will cancel out of any comparison across agents or comparison of an agent to the threshold. Hence we may as well simplify and ignore the λ .

In this case, the optimal mechanism is simply stated. If the incentive compatibility constraint does not bind, then $e_i = 1$ for every agent i and $\hat{v}_i = 1$. That is, every agent is asked for evidence and the good is allocated to the agent with the highest value. Hence incentive compatibility binds iff $c > 1/N$.

If the incentive compatibility constraint does bind, then $e_i = 1/(Nc)$ for all i . To see this, recall that the incentive compatibility constraint is that $E_{v_i} p_i(v_i) \geq c_i e_i$. In the symmetric case, all agents are equally likely to receive the good, so the left-hand side is $1/N$. Since this constraint binds and $c_i = c$ for all i , we see that $e_i = 1/(Nc)$ for all i . Finally, Lemma 5 implies that $\hat{v}$ is pinned down by

$$\sum_i e_i [1 - F_i(\hat{v}_i)] = 1 - \prod_i F_i(\hat{v}_i)$$

Using symmetry and rearranging, this is

$$e = \frac{1}{N} \left(\frac{1 - [F(\hat{v})]^N}{1 - F(\hat{v})} \right) = \frac{1}{N} \sum_{j=0}^{N-1} [F(\hat{v})]^j,$$

where the last equality comes from the formula for the sum of a geometric series. To see how this lines up with the dynamic mechanism, recall that we choose an order at random with, in the symmetric case, all orders equally likely.¹¹ This means that any given agent i has a $1/N$ chance of being first, $1/N$ chance of being second, etc. Think of the index j on the right-hand side as denoting how many agents are ahead of i in the selected order. If $j = 0$, then i is first and hence asked for evidence with certainty. If $j = 1$, there is one agent ahead of i and so i is asked for evidence iff this agent has a value below $\hat{v}$. Hence in this case, i is asked for evidence with probability $F(\hat{v})$. In general, there are j agents ahead of i with probability $1/N$ and i is asked for evidence in this situation with probability $[F(\hat{v})]^j$, the probability all these agents have values below $\hat{v}$. In short, using the value computed earlier for e , we see that $\hat{v}$ is defined by

$$\frac{1}{c} = \sum_{j=0}^{N-1} [F(\hat{v})]^j.$$

¹¹To be sure, there are other ways to generate the same evidence probabilities. For example, we could randomize uniformly over the orders $(1, 2, \dots, N)$, $(2, 3, \dots, N, 1)$, $(3, 4, \dots, N, 1, 2)$, $\dots$, $(N, 1, 2, \dots, N-1)$.

The comparative statics for the symmetric case are straightforward. If c increases, the left-hand side of the equation above falls, so $\hat{v}$ must fall. This is natural — if the cost of acquiring evidence goes up, agents must be promised a higher chance of getting the good to induce them to obtain evidence. This requires reducing the threshold. The evidence probability e also is reduced. Again, this fits with the reduction in the threshold. With a lower threshold, each agent is less likely to be reached and asked for evidence.

If the distribution F of values is shifted up in the sense of first-order stochastic dominance, then $F(\hat{v})$ is smaller at every point. Hence we must increase $\hat{v}$ to restore equality. Again, this is intuitive: if agents are more likely to have high values, then the principal can raise the threshold and be “pickier.” Note that e is unchanged since it is $1/(Nc)$. So the adjustment of $\hat{v}$ entirely eliminates what would otherwise be an improvement in an agent’s probability of getting the good from the improvement in F .

Finally, the effects of increasing the number of agents, N , is similarly straightforward to compute. If we increase N , then we must reduce $\hat{v}$ to restore equality in the equation above. So if there are more agents, the principal holds each to a *lower* standard, a perhaps unexpected conclusion. Intuitively, the principal does this because the increase in the number of agents reduces any one agent’s likelihood of getting the good if the threshold is unchanged. Thus each agent’s incentive to obtain evidence is reduced and must be restored by lowering the threshold. Increasing N also lowers the probability the agent is asked for evidence.

One might expect that the principal would be better off excluding some agents to keep the threshold higher. While it can be optimal for the principal to exclude some agents in asymmetric settings, this is never true in the symmetric case. Corollary 4 showed that if i is stronger than j and i is excluded, then j is excluded as well. In the symmetric case, all agents are stronger than all other agents. Hence if one of them is excluded in the optimal mechanism, all must be. But this can never be optimal since the principal could do better simply by giving the good to one of the agents without asking anyone for evidence.

To get some intuition for this, consider the expected number of agents asked for evidence. We know that if the principal asks, say, n agents for evidence, then each agent has an equal probability of being one of the n agents asked. Hence each agent has probability n/N of being asked for evidence in this case. So, overall, the expected probability that an agent is asked for evidence is $1/N$ times the expected number of agents asked. Since we know this probability is $1/(Nc)$, this says that the expected number of agents asked for evidence is $1/c$, independent of N . On the other hand, if the number of agents goes from N to $N + 1$, the probability the principal asks $N + 1$ agents for evidence goes from 0 to something strictly positive. Because the threshold falls, the probability the principal asks only one agent for evidence also increases. In other words,

as we increase the number of agents, the optimal mechanism changes so as to generate a mean-preserving spread in the number of agents asked. This change is valuable to the principal. It enables him to sample more agents when the draws are all low, though at the cost of sampling fewer when the draws are high.

The following theorem summarizes the results of this subsection:

Theorem 4. *When $c_i = c$ and $F_i = F$ for all i , the optimal mechanism gives the good to the agent with the highest v_i iff $c \leq 1/N$. Otherwise, the optimal mechanism has one tier and uses threshold $\hat{v}$ defined by*

$$\frac{1}{c} = \sum_{j=0}^{N-1} [F(\hat{v})]^j.$$

If F is shifted down in the sense of first-order stochastic dominance, if c is increased, or if N is increased, $\hat{v}$ falls.

4.3 Comparative Statics

While comparative statics are relatively straightforward in the symmetric case, this is not true in the asymmetric case, as we illustrate using Figure 2.

The figure is based on the case of two agents, both with $v_i \sim U[0, 1]$, with $c_2 = .62$ and c_1 varying between .5 and .75. Figure 2a shows how the indices of the two agents vary and consequently the probabilities they are asked for evidence. When c_1 is below the level shown by the first dotted line, 1's index is strictly larger than 2's, so the optimal mechanism starts by asking 1 for evidence. Consequently, 1 is asked for evidence with probability 1. In this range, 2 is only asked for evidence if 1's virtual value is below 2's index — that is, only if $v_1 < \hat{v}_2 + \lambda_2 - \lambda_1$. When c_1 is above the level shown by the third dotted line, 2's index is strictly larger and the roles are reversed. In between, the indices are the same, so the two agents are in the same tier. The probability that 1 is asked first for evidence declines with c_1 , so e_1 is decreasing and e_2 increasing in this range.

This last statement is seen more clearly in Figure 2b. The curve labeled q gives the probability 1 is asked for evidence first in the optimal mechanism. These curves give variables sufficient to compute the payoff of each agent. The curves $\hat{v}_2 + \lambda_2 - \lambda_1$ and $\hat{v}_1 + \lambda_1 - \lambda_2$ show what an agent's v_i is compared to if she is asked for evidence first, the former if $i = 1$, the latter if $i = 2$. The curves are dashed in the region where the other agent is first with probability 1 so that the curve is not relevant. The curve labeled $\lambda_2 - \lambda_1$ shows the relative advantage or disadvantage to an agent when both are asked for evidence as in this case, we compare $v_1 + \lambda_1$ to $v_2 + \lambda_2$.

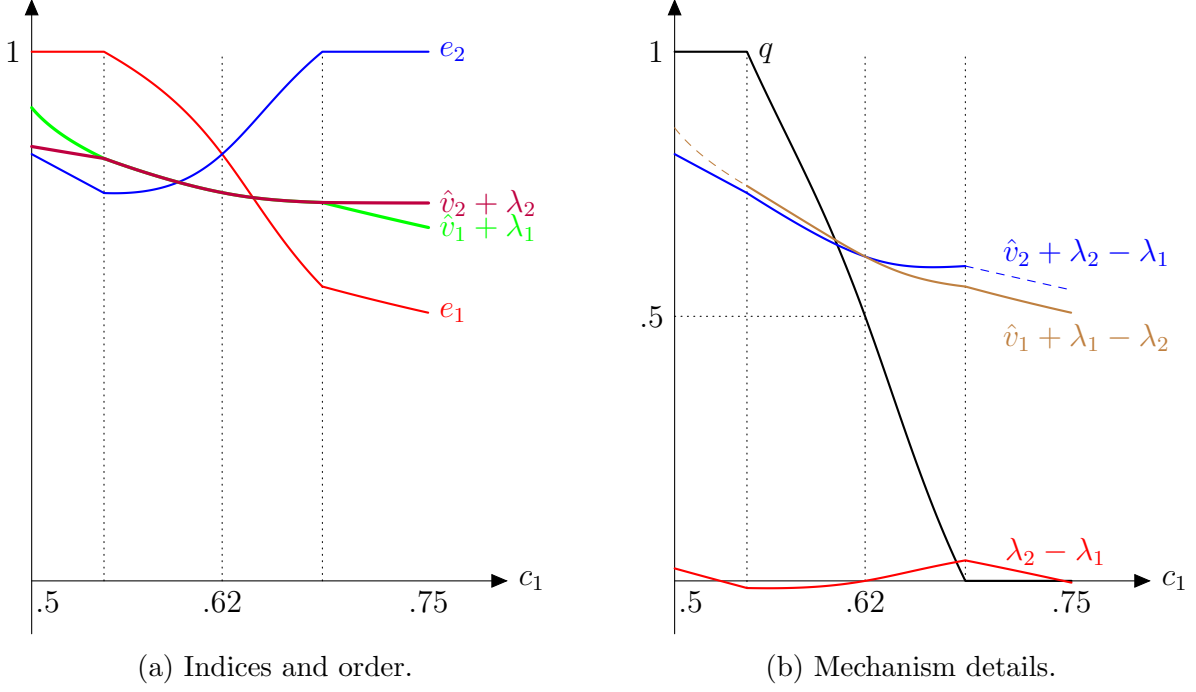


Figure 2: Illustrating the optimal mechanism for two agents, $v_i \sim U[0, 1]$, $c_2 = .62$, $c_1 \in [.5, .75]$.

To see the implications for comparative statics, first consider an increase in c_1 in the range where 1's index is strictly larger. If the increase is small, 1's index remains strictly larger, so e_1 and q both remain equal to 1. There are only two relevant variables left to manipulate, $\hat{v}_2 + \lambda_2 - \lambda_1$ and $\lambda_2 - \lambda_1$. The utility of each agent is decreasing in the first — 1 because she is more likely to have to compete with 2 and 2 because she has to compete with stronger 1 types — while 2's utility is increasing and 1's decreasing in the second. Hence the only way to restore incentive compatibility with equality for both agents is to lower both variables. Similar comments apply to the range where 2's index is strictly larger.

On the other hand, in between, we have more variables we can adjust. In this case, the optimal mechanism has a strictly interior probability of asking 1 for evidence first. Now the simple analysis above falls apart and a wide range of things can happen. This is simply because the principal now has another tool to compensate 1 for the increase in her costs. In addition to the kind of changes discussed above, the principal could now respond by lowering e_1 or, equivalently, lowering the probability that 1 is the first to be asked for evidence. The situation is even more complex in the general N -agent case. Because of this added flexibility, we have not found any comparative statics that hold globally.

5 Extensions

5.1 High Costs/Reserve Value

In this section, we discuss how the analysis changes if some agents have $c_i \geq 1$ or if the principal has a reservation value for the good.

Clearly, if some agent i has $c_i \geq 1$, the principal cannot get (useful) information from her. If $c_i > 1$, then i strictly prefers not getting the good to getting it and providing evidence. If $c_i = 1$, the principal can induce i to obtain evidence, but only by promising to give i the good for all v_i . Hence the principal cannot get information he can actually use from i .

As an aside, note that if we assume paying the cost is necessary to consume the good, then an agent with $c_i > 1$ does not want the good. In this case, such agents are irrelevant and the solution is the same as above with these agents removed.¹² The remainder of this section assumes paying the cost is not necessary for consuming the good.

In short, if there are agents with $c_i \geq 1$, the only role of such agents is as a kind of reservation value for the principal. In other words, we have assumed that the principal receives a payoff of 0 from keeping the good. If there is some agent i with $c_i \geq 1$, then the principal has a better outside option than keeping the good since he can give it to i for a “known” payoff of $E(v_i)$. Hence adding such an agent simply changes the principal’s reservation utility.

When there is a reservation value, say v_0 , for the principal, this has the same effect as adding an agent with a cost of 0 and a known value of v_0 , so that this “agent” has a Weitzman index of v_0 . This does not change the basic structure of the mechanism — we can still compute the Weitzman indices for the other agents in the same way, etc. While the general structure of the optimal mechanism is the same, the specific values of thresholds, etc., depend on v_0 in general.

Consider, for example, the symmetric case where all agents have distribution function F and cost c . If the reservation value v_0 is above the common Weitzman index, the principal takes the outside option and never considers the agents. Otherwise, the principal only takes the outside option if he checks all agents and all of them have virtual values below v_0 . Hence in the relevant range, the equations defining the common $\hat{v}$ and common

¹²The results in this subsection are the only ones in the paper which would change if we assumed that an agent cannot consume the good without paying the information cost c_i .

λ become

$$c \sum_{j=0}^{N-1} [F(\hat{v})]^j = 1 - [F(v_0 - \lambda)]^N$$

and

$$\lambda c = \int_{\hat{v}}^1 (v - \hat{v}) f(v) dv.$$

The first equation says that the total expected evidence costs of the agents add up to the total probability one of the agents receives the good, while the latter defines the Weitzman index. It is not hard to show that $\hat{v}$ is decreasing and λ increasing in v_0 in the relevant range where $v_0 \in (\lambda, \hat{v} + \lambda)$. Intuitively, because the principal may wish to take the outside option at the end, each agent must have a lower probability of being asked for evidence and so the principal has to be more lenient in the stopping decision.

5.2 Varying Value to Agent from Receiving Good

The use of Weitzman's results to characterize the structure of the optimal mechanism enables us to extend the analysis to the case where the agent's value of receiving the good varies and may be correlated with the value to the principal of giving it to her.

Suppose that if the value to the principal of giving the good to the agent is v_i , then the value to agent i of receiving it is $\varphi_i(v_i) > 0$. Assume $\varphi_i(\cdot)$ is continuous. We can normalize the agents' payoffs so that $E_{v_i} \varphi_i(v_i) = 1$ and assume, as above, that $c_i \in (0, 1)$ given this normalization. Then it is not hard to show that our result that every optimal mechanism satisfies no-free-lunch continues to hold. Now the Lagrangian takes the form

$$E_v \left[\sum_i P_i(v | d) v_i + \sum_i \lambda_i (P_i(v | d) \varphi_i(v_i) - e_i(d) c_i) \right]$$

or

$$E_v \left[\sum_i P_i(v | d) (v_i + \lambda_i \varphi_i(v_i)) - \sum_i \lambda_i c_i e_i(d) \right].$$

As above, we can view the λ 's as fixed parameters and characterize the solution to this maximization problem using the Weitzman solution for the case where the prize in box i is $v_i + \lambda_i \varphi_i(v_i)$. It is not hard to generalize the arguments above to show that there is an optimal mechanism which is a simple tiered threshold mechanism.

The function $\varphi_i(v_i)$ captures correlation between the value of the object to agent i and the value to the principal of giving her the object. For example, if we think of the principal as the dean of a college, the agents as departments, and the good as a job slot, it could be that the dean cares only about teaching, the departments only about research,

and that these are negatively correlated. If so, φ_i would be a decreasing function. In this case, we see the intuitive result that the need to incentivize agent i could imply that agent i is more likely to receive the good for some low values of v_i than for some higher values.

5.3 Costs

If we change the principal's objective so that he dislikes imposing costs on the agents, then with no other changes in the model, we lose the no-free-lunch property and hence the ability to appeal to Weitzman (1979). On the other hand, we can still use Weitzman's result and apply our analysis for a variation of our model.

To be specific, recall from the introduction that in some settings, it is natural to assume that the agent cannot consume the good without paying the cost. Consider this case and assume the principal's payoff if he allocates the good to agent i is v_i plus α_j times the utility of agent j , summed over the agents.

In this case, the principal can give the good to an agent without seeing evidence from her, but he factors in that the agent he gives it to will pay the cost. Hence it is as if the principal had to restrict attention to mechanisms satisfying no-free-lunch. In this case, given the λ_i 's, the Lagrangian reduces to the Weitzman problem where the prize in box i is $v_i + \lambda_i + \alpha_i$ and the cost of opening box i is $(\lambda_i + \alpha_i)c_i$. The incentive constraints are unaffected by this change, so the analysis is similar to the above. Again, there is a simple tiered threshold mechanism which is an optimal mechanism. Analogous comments apply if the principal internalizes the agent's value and cost differently so that, for example, the principal puts weight α_i on agent i 's cost and β_i on agent i 's utility from getting the good. Then the value of box i in the Weitzman problem is $v_i + \lambda_i + \beta_i$ and the cost is $(\lambda_i + \alpha_i)c_i$.

A Completion of Proof for Lemma 5

First, we show the implications of complementary slackness. Suppose there is a mechanism in $D^{**}(\lambda)$ satisfying complementary slackness which generates e . We show that $e_i = 1$ for any i with $\lambda_i = 0$.

So fix any i with $\lambda_i = 0$. Since $\hat{v}_i$ is defined by

$$\lambda_i c_i = \int_{\hat{v}_i}^1 (v - \hat{v}_i) f_i(v) dv,$$

$\lambda_i = 0$ implies $\hat{v}_i = 1$. Hence i 's index, $\hat{v}_i + \lambda_i$, equals 1.

Consider any $j \neq i$ for whom $\lambda_j > 0$ and $e_j > 0$. The same reasoning as above shows that we must have $\hat{v}_j < 1$. By complementary slackness, $\lambda_j > 0$ implies that j 's expected payoff in the mechanism must be zero. Hence $e_j[1 - F_j(\hat{v}_j) - c_j] \leq 0$ as agent j certainly gets the good if asked for evidence when her value is above $\hat{v}_j$ and might get the good even when her value is below $\hat{v}_j$. By assumption, $e_j > 0$, so $1 - F_j(\hat{v}_j) \leq c_j$.

To see the implication of this, note that j 's index satisfies

$$\hat{v}_j + \lambda_j = \hat{v}_j + \frac{1}{c_j} \int_{\hat{v}_j}^1 (v - \hat{v}_j) f_j(v) dv.$$

With this in mind, consider the function

$$\Phi_j(\hat{v}) \equiv \hat{v} + \frac{1}{c_j} \int_{\hat{v}}^1 (v - \hat{v}) f_j(v) dv.$$

Clearly,

$$\Phi_j'(\hat{v}) = 1 - \frac{1 - F_j(\hat{v})}{c_j}.$$

As shown above, $1 - F_j(\hat{v}_j) \leq c_j$, so $1 - F_j(\hat{v}) \leq c_j$ for all $\hat{v} \geq \hat{v}_j$, strictly if $\hat{v} > \hat{v}_j$. Hence $\hat{v}_j < 1$ implies

$$\hat{v}_j + \lambda_j < \Phi_j(1) = 1 = \hat{v}_i + \lambda_i.$$

Next, consider any j with $\lambda_j > 0$ and $e_j = 0$. For such j , we cannot have $\hat{v}_j + \lambda_j > 1$. If so, the j with the largest value of $\hat{v}_j + \lambda_j$ would be in the top tier, above any agent k with $\lambda_k = 0$ or $\lambda_k > 0$ and $e_k > 0$. But then we cannot have $e_j = 0$, a contradiction. Similarly, we cannot have $\hat{v}_j + \lambda_j = 1$. In this case, j and any agent k with $\lambda_k = 0$ are in the highest tier. Since $\hat{v}_k = 1$ for any k with $\lambda_k = 0$, all of these agents will be asked for evidence and the mechanism will continue to another agent with probability 1. Hence,

again, it would be impossible to have $e_j = 0$. Summarizing, if $\lambda_j > 0$, we must have $\hat{v}_j + \lambda_j < 1$.

Hence i is in the highest tier and any agent $j \neq i$ in the same tier also has $\lambda_j = 0$ and $\hat{v}_j = 1$. So all agents in this tier will be asked for evidence with probability 1, so $e_i = 1$.

The completion of the proof of Lemma 5 uses a result in Border (1991). For this section of the Appendix and the related Supplemental Appendix, we consider a different allocation problem with M agents where agent i has a *finite* set of types T_i . Types are independent across agents and μ_i is the distribution over T_i . We consider *allocations* $P = (P_1, \dots, P_N)$ with $P_i : T \rightarrow [0, 1]$ with $\sum_i P_i(t) \leq 1$ for all $t \in T$. Given P , $p = (p_1, \dots, p_N)$ denotes the *interim allocation generated by P* in the sense that

$$p_i(t_i) = \sum_{t_{-i} \in T_{-i}} \mu_{-i}(t_{-i}) P_i(t_i, t_{-i}).$$

A *hierarchical allocation* is an allocation P that can be constructed as follows. We have a ranking function R which maps $\cup_i T_i$ to $\{1, \dots, K\}$ for some K . We assume that for every $k < K$, there is exactly one i such that $R(t_i) = k$ for some $t_i \in T_i$. Note that this restriction does *not* apply to rank K — there may be no or many agents with types at rank K .

Then given a type profile $t = (t_1, \dots, t_N)$, either all agents have rank K or there is a unique i with $R(t_i) < R(t_j)$ for all $j \neq i$. If all agents have rank K , then $P_j(t) = 0$ for all j . If there is a unique i with $R(t_i) < R(t_j)$ for all $j \neq i$, then $P_i(t) = 1$. In other words, unless all agents are in the lowest rank, the agent who has the highest ranked type receives the good (where higher ranks have lower numbers).

Say that p is a *hierarchical interim allocation* if it is generated by a hierarchical allocation P . (The hierarchical interim allocations form a subset of the interim allocations.) The following is essentially¹³ Border's Lemma 6.1. For the reader's convenience, we include a proof of this result in the Supplemental Appendix.

Theorem 5. *Every interim allocation function p is a convex combination of hierarchical interim allocations.*

We also use the following finite type version of Border's theorem (Border (2007)):

¹³The differences are minor. Aside from differences in notation, the main change is that Border assumes iid types. However, as readers familiar with Border will recognize, the iid assumption is irrelevant to this result.

Theorem 6. *p is an interim allocation function if and only if for every collection $\hat{T}_i \subseteq T_i$ for $i = 1, \dots, M$, we have*

$$\sum_i \sum_{t_i \in \hat{T}_i} p_i(t_i) \mu_i(t_i) \leq 1 - \prod_i [1 - \mu_i(\hat{T}_i)].$$

To complete the proof of Lemma 5, fix $(e_1, \dots, e_N)$ satisfying equations (1) and (2) and with the property that $e_i = 1$ for all i with $\lambda_i = 0$. We show there exist randomizations over the orderings generating these evidence probabilities.

First, consider tier 1, $\mathcal{I}_1$. If there are any agents i with $\lambda_i = 0$, then tier 1 consists of the set of all such agents. In this case, since all are asked for evidence with probability 1, any randomization over the order in which to ask them will suffice.

So consider instead the case where there are no agents i with $\lambda_i = 0$, so we have $\hat{v}_i < 1$ for all i . The Border conditions (1) and (2) for tier 1 imply

$$\sum_{i \in \mathcal{I}_1} e_i [1 - F_i(\hat{v}_i)] = 1 - \prod_{i \in \mathcal{I}_1} F_i(\hat{v}_i) \quad (3)$$

and

$$\sum_{i \in \mathcal{I}} e_i [1 - F_i(\hat{v}_i)] \leq 1 - \prod_{i \in \mathcal{I}} F_i(\hat{v}_i), \quad \forall \mathcal{I} \subseteq \mathcal{I}_1. \quad (4)$$

We now construct an auxiliary allocation problem, the solution of which will provide the next step of the proof. Let the set of agents be $\mathcal{I}_1$. Each agent $i \in \mathcal{I}_1$ has two types, denoted ℓ_i and h_i , where $\mu_i(\ell_i) = F_i(\hat{v}_i)$. Note that $\hat{v}_i < 1$ implies $\mu_i(\ell_i) < 1$. Define functions $\hat{p}_i : T_i \rightarrow [0, 1]$ for $i \in \mathcal{I}_1$ by

$$\hat{p}_i(t_i) = \begin{cases} 0, & \text{if } t_i = \ell_i; \\ e_i, & \text{if } t_i = h_i. \end{cases}$$

Equations (3) and (4) and Theorem 6 imply that $\hat{p}$ is an interim allocation function.

By Theorem 5, $\hat{p}$ is a convex combination of hierarchical interim allocations. I.e., there are hierarchical interim allocations $q^1, \dots, q^S$ and weights $\alpha_s \in (0, 1)$ with $\sum_s \alpha_s = 1$ with $\hat{p} = \sum_s \alpha_s q^s$. Clearly, $\hat{p}_i(\ell_i) = 0$ for all i implies $q_i^s(\ell_i) = 0$ for all i and all s . By Theorem 6, we have

$$\sum_{i \in \mathcal{I}_1} q_i^s(h_i) [1 - F_i(\hat{v}_i)] \leq 1 - \prod_{i \in \mathcal{I}_1} F_i(\hat{v}_i), \quad \forall s.$$

But

$$\sum_s \alpha_s \left\{ \sum_{i \in \mathcal{I}_1} q_i^s(h_i) [1 - F_i(\hat{v}_i)] \right\} = \sum_{i \in \mathcal{I}_1} e_i [1 - F_i(\hat{v}_i)] = 1 - \prod_{i \in \mathcal{I}_1} F_i(\hat{v}_i).$$

Hence

$$\sum_{i \in \mathcal{I}_1} q_i^s(h_i)[1 - F_i(\hat{v}_i)] = 1 - \prod_{i \in \mathcal{I}_1} F_i(\hat{v}_i), \quad \forall s.$$

The left-hand side is the probability that some agent i is type h_i and receives the good, while the right-hand side is the probability at least one agent is type h_i . So this equality says that for every s , if at least one agent is type h_i , the good is allocated to such an agent. Because $\mu_i(\ell_i) < 1$ for all i , this event has strictly positive probability.

Given this, consider any q^s . Since q^s is a hierarchical interim allocation, there is a hierarchical allocation, Q^s , and a ranking function, R^s , associated with it. From the above, we know that for any type profile such that some i is type h_i , the good is allocated to such an i . Because the allocation is hierarchical, there is a unique such i who gets the good with probability 1. Consider the allocation on type profile $(h_1, \dots, h_{\#\mathcal{I}_1})$. Whichever agent i receives the good on this profile must have $R^s(h_i) = 1$. Let i_1 denote this agent.

Consider the profile of types where agent i_1 is type ℓ_{i_1} and every other agent i is type h_i . Again, there must be an agent, say i_2 , who receives the good with probability 1 and hence we have $R^s(h_{i_2}) = 2$. Continuing this way, we construct the ranking R^s which orders the h_i types of all agents. Define an ordering over $i \in \mathcal{I}_1$, $\succ^s$, by $i \prec^s j$ iff $R^s(h_i) < R^s(h_j)$.

By construction,

$$e_i = \hat{p}_i(h_i) = \sum_s \alpha^s q_i^s(h_i) = \sum_s \alpha^s \prod_{j \prec^s i} \mu_j(\ell_j) = \sum_s \alpha^s \prod_{j \prec^s i} F_j(\hat{v}_j).$$

So the randomization over orderings of $\mathcal{I}_1$ given by $O_1(\succ^s) = \alpha^s$ generates the e_i 's for $\mathcal{I}_1$.

Next consider tier 2, $\mathcal{I}_2$. By equations (1) and (2), we know that the evidence probabilities for agents in this tier satisfy

$$\sum_{i \in \mathcal{I}_2} e_i [1 - F_i(\hat{v}_i)] = \left[\prod_{i \in \mathcal{I}_1} F_i(v_2^* - \lambda_i) \right] \left[1 - \prod_{i \in \mathcal{I}_2} F_i(\hat{v}_i) \right] \quad (5)$$

and

$$\sum_{i \in \mathcal{I}} e_i [1 - F_i(\hat{v}_i)] \leq \left[\prod_{i \in \mathcal{I}_1} F_i(v_2^* - \lambda_i) \right] \left[1 - \prod_{i \in \mathcal{I}} F_i(\hat{v}_i) \right], \quad \forall \mathcal{I} \subseteq \mathcal{I}_2. \quad (6)$$

For $i \in \mathcal{I}_2$, let

$$\hat{e}_i = \frac{e_i}{\prod_{j \in \mathcal{I}_1} F_j(v_2^* - \lambda_j)}.$$

Then equations (3) and (4) hold for tier 2 and the $\hat{e}$'s. That is, equations (5) and (6) can be rewritten as

$$\sum_{i \in \mathcal{I}_2} \hat{e}_i [1 - F_i(\hat{v}_i)] = 1 - \prod_{i \in \mathcal{I}_2} F_i(\hat{v}_i)$$

$$\sum_{i \in \mathcal{I}} \hat{e}_i [1 - F_i(\hat{v}_i)] \leq 1 - \prod_{i \in \mathcal{I}} F_i(\hat{v}_i), \quad \forall \mathcal{I} \subseteq \mathcal{I}_2.$$

Hence the same argument as above shows that we can construct a probability distribution O_2 over orderings $\succ^s$ over $\mathcal{I}_2$ such that for all $i \in \mathcal{I}_2$,

$$\hat{e}_i = \sum_s O_2(\succ^s) \prod_{j \prec^s i} F_j(\hat{v}_j)$$

or, equivalently,

$$e_i = \left[\prod_{j \in \mathcal{I}_1} F_j(v_2^* - \lambda_j) \right] \sum_s O_2(\succ^s) \prod_{j \prec^s i} F_j(\hat{v}_j).$$

Iterating this argument for the remaining tiers completes the proof. ■

B Proof of Theorem 3

Lemma 6. *Suppose i is stronger than j . In any optimal mechanism with $e_j > 0$, we have $e_i > 0$.*

Proof. The proof is by contradiction. So suppose i is stronger than j , but we have an optimal mechanism d^1 with $e_i(d^1) = e_i^1 = 0$ and $e_j(d^1) = e_j^1 > 0$. By no-free-lunch, we have $P_i(v \mid d^1) = 0$ for (almost) all v . Given this, we must have at least three agents. Otherwise, a best outcome of this form is to simply give the good to j with probability 1. Essentially the same argument as the proof of Lemma 1 gives a contradiction to this being optimal.

We write $\hat{v}_k^1, \lambda_k^1$, etc., to denote the relevant variables for this outcome. We write P^1 for $P(d^1)$ and e^1 for $e(d^1)$.

We construct a second mechanism d^2 as follows. Define a function $\varphi : V_i \rightarrow V_j$ by $\varphi(v_i) = F_j^{-1}(F_i(v_i))$. Note that the distribution of $\varphi(v_i)$ is the same as the distribution of v_j . That is, for any z , we have $\Pr[\varphi(v_i) \leq z] = F_j(z)$.

Define d^2 to be the same as d^1 except as follows. First, on any history where d^1 asks j for evidence with positive probability, d^2 asks i instead with this same probability. Second, if i is asked for evidence and proves value v_i , d^2 treats this the same way mechanism

d^1 treats proof by j of value $\varphi(v_i)$. In the Supplemental Appendix, we show that d^2 is incentive compatible and weakly increases the principal's expected payoff. Since d^1 is optimal, this implies d^2 must yield the same payoff as d^1 .

Let $d^* = (1/2)d^1 + (1/2)d^2$, the 50–50 randomization between these mechanisms. We write P^* for $P(d^*)$, etc. Since both d^1 and d^2 are optimal, so is d^* . Since every optimal mechanism is a generalized tiered threshold mechanism, so is d^* .

Note that $e_i^* = e_j^* = (1/2)e_j^1 > 0$, so i and j must be in the same tier. That is, we must have $\hat{v}_i^* + \lambda_i^* = \hat{v}_j^* + \lambda_j^*$.

On any history (where we include the randomization over the order if i and j are in a tier with one or more other agents) where d^1 would have asked j for evidence, the mechanism d^* randomizes 50–50 between asking i or asking j . Whichever is chosen, the mechanism *never* continues to the other. For simplicity, consider the case where the mechanism asks i for evidence and never continues to j (the other case is symmetric). The only way this can be consistent with a generalized tiered threshold mechanism is if either $\hat{v}_i^* \leq 0$ or for some k in the same tier as i and j we have $e_k > 0$ and $\hat{v}_k^* \leq 0$. Otherwise, there is a strictly positive probability that $v_i \in [0, \hat{v}_i^*)$ and $v_k \in [0, \hat{v}_k^*)$ for every $k \neq j$ with $e_k > 0$ who has not already been asked. In this case, we *must* continue to j , a contradiction.

So consider the agent k (where k may be equal to i) with $\hat{v}_k^* \leq 0$. Then agent k receives the good iff she is asked for evidence. Hence her expected utility in the mechanism is $e_k^* - c_k e_k^* = (1 - c_k)e_k^* > 0$, so k 's incentive constraint is not binding. This implies $\lambda_k^* = 0$, so $\hat{v}_k^*$ satisfies

$$0 = \int_{\hat{v}_k^*}^1 (v - \hat{v}_k^*) f_k(v) dv$$

requiring $\hat{v}_k^* = 1$, a contradiction. ■

We prove Theorem 3 by contradiction. So fix an optimal mechanism d . Suppose i is stronger than j but $e_i < e_j$. By Lemma 5, $e_i < e_j \leq 1$ implies $\lambda_i > 0$. Lemma 6 rules out the possibility that $0 = e_i < e_j$, so we must have $0 < e_i < e_j$.

Lemma 7. *If i is stronger than j and $0 < e_i < e_j$, then $\hat{v}_j + \lambda_j \geq \hat{v}_i + \lambda_i$, $\hat{v}_j \geq \hat{v}_i$, and $\lambda_j \leq \lambda_i$.*

Proof of Lemma. Clearly, if $e_j > e_i$, we cannot have j in a lower tier than i . Any agent is only asked for evidence after all agents in higher tiers are asked, so $e_j > e_i$ implies j is in a weakly higher tier than i . Hence $\hat{v}_j + \lambda_j \geq \hat{v}_i + \lambda_i$, establishing the first claim.

Define

$$\Lambda_i(\hat{v}) = \frac{1}{c_i} \int_{\hat{v}}^1 (v - \hat{v}) f_i(v) dv \quad (7)$$

The index $\hat{v}_i + \lambda_i$ satisfies $\hat{v}_i + \lambda_i = \hat{v}_i + \Lambda_i(\hat{v}_i)$. Also, in the proof of Lemma 5, we defined $\Phi_i(\hat{v}) = \hat{v} + \Lambda_i(\hat{v})$ and showed that $\Phi_i(\hat{v})$ is strictly increasing in $\hat{v}$ for all $\hat{v} > \hat{v}_i$ if $e_i > 0$.

Next, we show that if i is stronger than j , then for all $\hat{v}$, we have $\Lambda_i(\hat{v}) \geq \Lambda_j(\hat{v})$. This holds with equality if $c_i = c_j$ and $F_i = F_j$ since the functions are then the same. If $F_i = F_j$ but $c_i < c_j$, the Λ_i function must be larger at every $\hat{v} < 1$ as it's defined by dividing by a strictly smaller number. Alternatively, suppose $c_i = c_j$ but F_i FOSD F_j . Because the function $\max\{0, v - \hat{v}\}$ is increasing in v , F_i FOSD F_j implies

$$\int_{\hat{v}}^1 (v - \hat{v}) f_i(v) dv \geq \int_{\hat{v}}^1 (v - \hat{v}) f_j(v) dv,$$

again implying $\Lambda_i(\hat{v}) \geq \Lambda_j(\hat{v})$.

Suppose, contrary to our claim, that $\hat{v}_i > \hat{v}_j$. Note that $e_j > 0$, so $\hat{v} + \Lambda_j(\hat{v})$ is strictly increasing in $\hat{v}$ for all $\hat{v} \in (\hat{v}_j, \hat{v}_i)$, implying $\hat{v}_j + \Lambda_j(\hat{v}_j) < \hat{v}_i + \Lambda_j(\hat{v}_i)$. Then i stronger than j implies $\hat{v}_j + \Lambda_j(\hat{v}_j) < \hat{v}_i + \Lambda_j(\hat{v}_i) \leq \hat{v}_i + \Lambda_i(\hat{v}_i)$. From equation (7), we have $\hat{v}_j + \lambda_j = \hat{v}_j + \Lambda_j(\hat{v}_j) < \hat{v}_i + \Lambda_j(\hat{v}_i) \leq \hat{v}_i + \Lambda_i(\hat{v}_i) = \hat{v}_i + \lambda_i$. This contradicts j being in a weakly higher tier than i , so $\hat{v}_j \geq \hat{v}_i$, proving the second claim.

Finally, the $\Lambda_n(\cdot)$ functions are decreasing, so (7) implies $\lambda_i = \Lambda_i(\hat{v}_i) \geq \Lambda_i(\hat{v}_j)$. By i stronger than j , we have $\Lambda_i(\hat{v}_j) \geq \Lambda_j(\hat{v}_j) = \lambda_j$. Hence $\lambda_i \geq \lambda_j$, showing the third claim. ■

We derive a contradiction from this lemma and $e_i, \lambda_i > 0$. From the proof of Lemma 4,

$$p_i(v_i | d) = \begin{cases} e_i, & \text{if } v_i \geq \hat{v}_i; \\ \prod_{j \neq i | \hat{v}_j + \lambda_j \geq v_i + \lambda_i} F_j(v_i + \lambda_i - \lambda_j), & \text{otherwise,} \end{cases}$$

so we can write the expected payoff to any agent k as

$$U_k = e_k[1 - F_k(\hat{v}_k)] - e_k c_k + \int_0^{\hat{v}_k} \left[\prod_{j \neq k | \hat{v}_j + \lambda_j \geq v_k + \lambda_k} F_j(v + \lambda_k - \lambda_j) \right] f_k(v) dv.$$

Because $\lambda_i > 0$, we must have $U_i = 0 \leq U_j$.

For $v_j < \hat{v}_j$, we have

$$p_j(v_j | d) = \prod_{k \neq j | \hat{v}_k + \lambda_k \geq v_j + \lambda_j} F_k(v_j + \lambda_j - \lambda_k).$$

Let $w = v_j + \lambda_j$. Then for $w \in [\lambda_i, \hat{v}_j + \lambda_j)$, we have

$$\begin{aligned}
p_j(w - \lambda_j \mid d) &= \prod_{k \neq j \mid \hat{v}_k + \lambda_k \geq w} F_k(w - \lambda_k) \\
&= F_i(w - \lambda_i) \prod_{k \neq i, j \mid \hat{v}_k + \lambda_k \geq w} F_k(w - \lambda_k) \\
&\leq F_j(w - \lambda_j) \prod_{k \neq i, j \mid \hat{v}_k + \lambda_k \geq w} F_k(w - \lambda_k) = \prod_{k \neq i \mid \hat{v}_k + \lambda_k \geq w} F_k(w - \lambda_k) \\
&= p_i(w - \lambda_i \mid d), \quad \forall w - \lambda_i < \hat{v}_i
\end{aligned}$$

where the inequality in the third line comes from F_i FOSD F_j and $\lambda_i \geq \lambda_j$. To understand the last line, note that $w - \lambda_i < \hat{v}_i$ means that $v_i = w_i - \lambda_i < \hat{v}_i$, so the formula given in the proof of Lemma 4 for $p_i(v_i \mid d)$ applies.

Summarizing, we have

$$p_j(v_j \mid d) \leq p_i(v_j + \lambda_j - \lambda_i \mid d), \quad \forall v_j < \hat{v}_i + \lambda_i - \lambda_j.$$

No agent gets the good without being asked for evidence, so $p_i(v'_i \mid d) \leq e_i$ for all v'_i and

$$p_j(v_j \mid d) \leq e_i, \quad \forall v_j < \hat{v}_i + \lambda_i - \lambda_j \leq \hat{v}_j. \quad (8)$$

Next, $p_i(v_i \mid d)$ increasing in v_i and F_i FOSD F_j implies

$$\int_0^1 p_i(v_i \mid d) f_i(v_i) dv_i \geq \int_0^1 p_i(v_i \mid d) f_j(v_i) dv_i.$$

Rewriting the right-hand side, we have

$$\int_0^1 p_i(v_i \mid d) f_j(v_i) dv_i = \int_0^{\hat{v}_i} \left[\prod_{k \neq i \mid \hat{v}_k + \lambda_k \geq v_i + \lambda_i} F_k(v_i + \lambda_i - \lambda_k) \right] f_j(v_i) dv_i + [1 - F_j(\hat{v}_i)] e_i.$$

Change variables in the integral on the right-side side by defining $w = v_i + \lambda_i$, so it is

$$\int_{\lambda_i}^{\hat{v}_i + \lambda_i} \left[\prod_{k \neq i \mid \hat{v}_k + \lambda_k \geq w} F_k(w - \lambda_k) \right] f_j(w - \lambda_i) dw.$$

The same reasoning as above shows that this is weakly larger than what we get if we take the product over $k \neq j$ instead of $k \neq i$ — that is, is weakly larger than

$$\int_{\lambda_i}^{\hat{v}_i + \lambda_i} \left[\prod_{k \neq j \mid \hat{v}_k + \lambda_k \geq w} F_k(w - \lambda_k) \right] f_j(w - \lambda_i) dw.$$

Change variables in the integral again, replacing w with $v_j + \lambda_i$ to write this as

$$\int_0^{\hat{v}_i} \left[\prod_{k \neq j | \hat{v}_k + \lambda_k \geq v_j + \lambda_i} F_k(v_j + \lambda_i - \lambda_k) \right] f_j(v_j) dv_j.$$

The fact that $\lambda_i \geq \lambda_j$ implies $F_k(v_j + \lambda_i - \lambda_k) \geq F_k(v_j + \lambda_j - \lambda_k)$ for all k and v_j . Also, if we change the index on the product from $k \neq j$ such that $\hat{v}_k + \lambda_k \geq v_j + \lambda_i$ to $k \neq j$ such that $\hat{v}_k + \lambda_k \geq v_j + \lambda_j$, $\lambda_i \geq \lambda_j$ implies we will be taking the product over weakly more k 's. Since each term in the product is weakly less than 1, this must reduce the product. Hence the expression above is weakly larger than

$$\int_0^{\hat{v}_i} \left[\prod_{k \neq j | \hat{v}_k + \lambda_k \geq v_j + \lambda_j} F_k(v_j + \lambda_j - \lambda_k) \right] f_j(v_j) dv_j = \int_0^{\hat{v}_i} p_j(v_j | d) f_j(v_j) dv_j.$$

The fact that i 's incentive constraint binds implies

$$c_i e_i = \int_0^1 p_i(v_i | d) f_i(v_i) dv_i \geq \int_0^{\hat{v}_i} p_j(v_j | d) f_j(v_j) dv_j + [1 - F_j(\hat{v}_i)] e_i,$$

while j 's incentive constraint implies

$$\int_0^{\hat{v}_j} p_j(v_j | d) f_j(v_j) dv_j + (1 - F_j(\hat{v}_j)) e_j \geq c_j e_j.$$

For $v_j \in [\hat{v}_i, \hat{v}_i + \lambda_i - \lambda_j]$, equation (8) implies $p_j(v_j | d) \leq e_i$. For $v_j \in [\hat{v}_i + \lambda_i - \lambda_j, \hat{v}_j]$, we have $p_j(v_j | d) \leq e_j$. Hence

$$\int_0^{\hat{v}_i} p_j(v_j | d) f_j(v_j) dv_j + [F_j(\hat{v}_i + \lambda_i - \lambda_j) - F_j(\hat{v}_i)] e_i + [1 - F_j(\hat{v}_i + \lambda_i - \lambda_j)] e_j \geq c_j e_j.$$

Summarizing, we have

$$\begin{aligned} c_i e_i &\geq \int_0^{\hat{v}_i} p_j(v_j | d) f_j(v_j) dv_j + [1 - F_j(\hat{v}_i)] e_i \\ &\geq [1 - F_j(\hat{v}_i)] e_i + c_j e_j - [F_j(\hat{v}_i + \lambda_i - \lambda_j) - F_j(\hat{v}_i)] e_i - [1 - F_j(\hat{v}_i + \lambda_i - \lambda_j)] e_j \\ &= c_j e_j + [1 - F_j(\hat{v}_i + \lambda_i - \lambda_j)] e_i - [1 - F_j(\hat{v}_i + \lambda_i - \lambda_j)] e_j. \end{aligned}$$

So

$$[c_i - (1 - F_j(\hat{v}_i + \lambda_i - \lambda_j))] e_i \geq [c_j - (1 - F_j(\hat{v}_i + \lambda_i - \lambda_j))] e_j. \quad (9)$$

Recall that

$$U_i = 0 = (1 - F_i(\hat{v}_i) - c_i) e_i + \int_0^{\hat{v}_i} \left[\prod_{k \neq i | \hat{v}_k + \lambda_k \geq v_i + \lambda_i} F_k(v_i + \lambda_i - \lambda_k) \right] f_i(v_i) dv_i.$$

Obviously, the integral is non-negative as it is a probability. We now show that $e_i > 0$ implies that the integral must be strictly positive. Recall that the mechanism does not ask i for evidence until every agent k with $\hat{v}_k + \lambda_k > \hat{v}_i + \lambda_i$ has already been asked for evidence and has been found to have a virtual value strictly below $\hat{v}_i + \lambda_i$. Hence the fact that $e_i > 0$ implies that this event must have positive probability.

Furthermore, any other agent k in the same tier as i must have a positive probability of having a virtual value below the tier threshold. If this were not true, we must have $\hat{v}_k \leq 0$. But if $\hat{v}_k \leq 0$, agent k receives the good if and only if she is asked for evidence, so her probability of getting the good ex ante is e_k . But then her expected payoff is $e_k(1 - c_k) > 0$ as we assume $c_k < 1$ for all k . This implies that $\lambda_k = 0$, which by the Weitzman formula implies $\hat{v}_k = 1$, not zero, a contradiction.

Hence $U_i = 0$ implies

$$0 < \int_0^{\hat{v}_i} \left[\prod_{k \neq i | \hat{v}_k + \lambda_k \geq v_i + \lambda_i} F_k(v_i + \lambda_i - \lambda_k) \right] f_i(v_i) dv_i = [c_i - (1 - F_i(\hat{v}_i))]e_i.$$

So we have

$$c_i > 1 - F_i(\hat{v}_i) \geq 1 - F_i(\hat{v}_i + \lambda_i - \lambda_j) \geq 1 - F_j(\hat{v}_i + \lambda_i - \lambda_j),$$

where the second inequality is implied by $\lambda_i \geq \lambda_j$ and the second from F_i FOSD F_j .

Hence $c_i - [1 - F_j(\hat{v}_i + \lambda_i - \lambda_j)] > 0$, so $e_j > e_i$ implies

$$[c_i - (1 - F_j(\hat{v}_i + \lambda_i - \lambda_j))]e_i < [c_i - (1 - F_j(\hat{v}_i + \lambda_i - \lambda_j))]e_j \leq [c_j - (1 - F_j(\hat{v}_i + \lambda_i - \lambda_j))]e_j,$$

where the second inequality follows from $c_j \geq c_i$. But this contradicts equation (9). ■

References

- [1] Ball, I., and D. Kattwinkel, “Probabilistic Verification in Mechanism Design,” *Theoretical Economics*, **20**, November 2025, 1247–1284.
- [2] Ben-Porath, E., E. Dekel, and B. Lipman, “Mechanisms with Evidence: Commitment and Robustness,” *Econometrica*, **87**, March 2019, 529–566.
- [3] Ben-Porath, E., E. Dekel, and B. Lipman, “Mechanism Design for Acquisition of/Stochastic Evidence,” *Theoretical Economics*, **21**, January 2026, 99–131.
- [4] Bergemann, D., and J. Välimäki, “Information in Mechanism Design,” in R. Blundell, W. Newey, and T. Persson, eds., *Advances in Economics and Econometrics: Theory and Applications, Ninth World Congress*, Cambridge University Press, 2006, 186–221.
- [5] Border, K., “Implementation of Reduced Form Auctions: A Geometric Approach,” *Econometrica*, **59**, July 1991, 1175–1187.
- [6] Border, K., “Reduced Form Auctions Revisited,” *Economic Theory*, **31**, April 2007, 167–181.
- [7] Bull, J., and J. Watson, “Hard Evidence and Mechanism Design,” *Games and Economic Behavior*, **58**, January 2007, 75–93.
- [8] Che, Y.-K., J. Kim, and K. Mierendorff, “Generalized Reduced-Form Auctions: A Network-Flow Approach,” *Econometrica*, **81**, November 2013, 2487–2520.
- [9] Crémer, J., Y. Spiegel, and C. Zheng, “Auctions with Costly Information Acquisition,” *Economic Theory*, **38**, January 2009, 41–72.
- [10] Deneckere, R., and S. Severinov, “Mechanism Design with Partial State Verifiability,” *Games and Economic Behavior*, **64**, November 2008, 487–513.
- [11] Doval, L., “Whether or Not to Open Pandora’s Box,” *Journal of Economic Theory*, **175**, 2018, 127–158.
- [12] Gershkov, A., and B. Szentes, “Optimal Voting Schemes with Costly Information Acquisition,” *Journal of Economic Theory*, **144**, January 2009, 36–68.
- [13] Glazer, J., and A. Rubinstein, “On Optimal Rules of Persuasion,” *Econometrica*, **72**, November 2004, 1715–1736.
- [14] Glazer, J., and A. Rubinstein, “A Study in the Pragmatics of Persuasion: A Game Theoretical Approach,” *Theoretical Economics*, **1**, December 2006, 395–410.

- [15] Green, J., and J.-J. Laffont, “Partially Verifiable Information and Mechanism Design,” *Review of Economic Studies*, **53**, July 1986, 447–456.
- [16] Hart, S., I. Kremer, and M. Perry, “Evidence Games: Truth and Commitment,” *American Economic Review*, **107**, March 2017, 690–713.
- [17] Khalfan, N., and R. Vohra, “Sequential Information Acquisition and Optimal Search,” Monash University working paper, October 2024.
- [18] Mierendorff, K., “Asymmetric Reduced Form Auctions,” *Economics Letters*, **110**, January 2011, 41–44.
- [19] Myerson, R., “Multistage Games with Communication,” *Econometrica*, **54**, March 1986, 323–358.
- [20] Weitzman, M., “Optimal Search for the Best Alternative,” *Econometrica*, **47**, May 1979, 641–654.