

# Randomization and the Robustness of Linear Contracts\*

Ashwin Kambhampati<sup>†</sup>    Bo Peng<sup>‡</sup>    Zhihao Gavin Tang<sup>§</sup>  
Juuso Toikka<sup>¶</sup>    Rakesh Vohra<sup>||</sup>

July 10, 2026

## Abstract

We consider contract design by an uncertainty-averse principal who does not know the production technology available to an agent. For such a principal, randomizing the choice of contract provides a hedge against uncertainty. We show that it is optimal to only use linear contracts and randomize the principal's share according to a log-uniform distribution. Thus, an optimal response to uncertainty may generate rich heterogeneity in contracts offered to observationally identical agents. Our saddle-point characterization of the optimal contract implies that the gain from randomization can be arbitrarily large; optimal randomization does not require commitment; and that asking the agent to report the technology cannot improve the principal's payoff.

## 1 Introduction

Recent years have witnessed a surge in contract-theoretic research on the optimal design of incentives when the designer does not know all the details of the contracting environment, or is otherwise reluctant to commit to a single probabilistic model of the situation. For moral hazard problems, the predominant approach is to assume that the designer evaluates each contract according to its payoff guarantee, i.e., its worst-case payoff across a class of

---

\*This paper merges two independent works, [Kambhampati, Toikka, and Vohra \(2024\)](#) and [Peng and Tang \(2024\)](#).

<sup>†</sup>Department of Economics, United States Naval Academy. Email: [kambhamp@usna.edu](mailto:kambhamp@usna.edu).

<sup>‡</sup>Key Laboratory of Interdisciplinary Research of Computation and Economics, Shanghai University of Finance and Economics. Email: [ahqspb@163.sufe.edu.cn](mailto:ahqspb@163.sufe.edu.cn).

<sup>§</sup>Key Laboratory of Interdisciplinary Research of Computation and Economics, Shanghai University of Finance and Economics. Email: [tang.zhihao@mail.shufe.edu.cn](mailto:tang.zhihao@mail.shufe.edu.cn).

<sup>¶</sup>The Wharton School, University of Pennsylvania. Email: [toikka@wharton.upenn.edu](mailto:toikka@wharton.upenn.edu).

<sup>||</sup>Department of Economics & Department of Electrical and Systems Engineering, University of Pennsylvania. Email: [rvohra@seas.upenn.edu](mailto:rvohra@seas.upenn.edu).

environments. This approach has been successful in providing a possible foundation for linear contracts, a commonly observed contract form—see [Carroll \(2015\)](#) and subsequent work by, e.g., [Dai and Toikka \(2022\)](#), [Walton and Carroll \(2022\)](#), [Carroll and Bolte \(2023\)](#), [Marku, Ocampo Diaz, and Tondji \(2024\)](#), and [Vairo \(2025\)](#). Moreover, it does so without the need to resort to specific assumptions about how actions map to outputs and costs, a challenge for the Bayesian contracting literature.

The robust analysis of moral hazard problems, however, typically restricts attention to deterministic contracts. This restriction is consequential. The robust contracting problem naturally maps to the setting of [Anscombe and Aumann \(1963\)](#), with the unknown aspects of the environment corresponding to the subjectively uncertain state of the world and objective risk arising from the stochastic mapping of effort to output in any given environment. Any contract together with the agent’s best response then induces an act, i.e., a mapping from states to profit distributions. Viewed this way, the key axiom underlying [Gilboa and Schmeidler’s \(1989\)](#) max-min representation of uncertainty aversion is precisely the assumption that randomizing the choice of contract provides a hedge against uncertainty.<sup>1</sup> Excluding randomization, therefore, amounts to a no-hedging restriction at odds with the preferences justifying the max-min approach. Note that no such restriction is typically imposed in the recent literature on robust mechanism design—see, e.g., [Carrasco, Luz, Kos, Messner, Monteiro, and Moreira \(2018\)](#), [Che and Zhong \(2024\)](#), [Deb and Roesler \(2024\)](#), and [Brooks and Du \(2021\)](#) who find randomized prices and mechanisms to be robustly optimal in single- and multi-dimensional screening and auction environments.

While the axiom of uncertainty aversion only requires a weak preference for randomization, [Kambhampati \(2023\)](#) shows that the preference is strict in the canonical framework of [Carroll \(2015\)](#): randomizing over two linear contracts strictly improves the principal’s payoff relative to any deterministic contract. However, [Kambhampati \(2023\)](#) does not show that linear contracts are optimal among random contracts. Thus, the result raises a foundational question for the literature: Are robustly optimal contracts still linear or, more broadly, simple, once the principal is allowed to hedge uncertainty? In other words, does the robust optimality of interest alignment established in the deterministic case still hold? The answer is not obvious because, as we elaborate below, allowing for random contracts renders the principal’s problem equivalent to designing an optimal information structure for a multi-dimensional mechanism design problem, which itself is well known to be difficult to solve in general (see, e.g., the recent survey by [Rochet, 2024](#)).

In this paper, we study robustly optimal incentive provision by an uncertainty-averse

---

<sup>1</sup>The axiom states that when indifferent over two acts, the principal (weakly) prefers any randomization between them to choosing either of them deterministically. See [Appendix C](#) for a formal treatment.

principal who can hedge against uncertainty by randomizing the choice of contract. For concreteness, we cast the analysis in the context of [Carroll’s \(2015\)](#) robust moral hazard problem. In the model, the principal knows only some actions the agent can take and evaluates contracts based on their guaranteed payoff over all possible action sets containing the known actions. The agent is protected by limited liability and both parties are risk-neutral. The principal chooses a randomization over contracts and offers the realized contract to the agent.

We show that the alignment of the principal’s and the agent’s interests achieved by linear contracts provides a robustly optimal way to incentivize the agent, whether the contract is deterministic or random. Specifically, any (finitely supported) random contract can be improved by linearizing each contract in its support while keeping the probabilities fixed. Thus, the problem of hedging against uncertainty can be decoupled from the problem of choosing the form of the contract given to the agent. The model therefore retains the prediction that the agent need only ever be offered a linear contract.

We then show that there exists an optimal random linear contract in which the principal’s share is drawn from a log-uniform distribution (i.e., the logarithm of the principal’s share has a uniform distribution). Thus, the optimal response to uncertainty can generate rich heterogeneity in contracts offered to observationally identical agents once we allow for hedging. Heterogeneity of course does not arise in the deterministic case and the richness contrasts with the binary improvement contract in [Kambhampati \(2023\)](#).

To establish these results, we formulate the guarantee of an arbitrary random contract as a linear program for adversarial Nature. This problem is in general a multi-dimensional mechanism design problem, where the agent’s type is the realized contract. Nature designs for each type an action subject to usual incentive compatibility constraints and a type-dependent participation constraint stemming from the presence of actions known to the principal. The random contract corresponds to the type distribution, and so the principal’s contract design problem amounts to designing an optimal information structure in a multi-dimensional mechanism design problem.

To prove that linear contracts remain optimal, we use the dual to Nature’s problem to identify a random linear improvement contract. Applied to the deterministic case, the approach yields a short proof of [Carroll’s \(2015\)](#) result, with linear programming duality replacing his separating hyperplane argument. The result provides a crucial simplification towards solving for an optimal random contract as Nature’s mechanism design problem becomes one-dimensional when attention is restricted to random linear contracts.

To find an optimal contract, we then consider randomization over a finite grid of linear contracts, identified by their slopes. This problem is made tractable by turning the principal’s

max-min problem into a max-max problem by again taking the dual of Nature’s minimization problem. The max-max problem is, of course, simply a maximization problem and can be solved for analytically. As the grid becomes dense in the unit interval, the solution converges to a contract where the principal’s share has log-uniform distribution; this contract is a natural guess for an optimal random contract. We verify its optimality by constructing a saddle point in the zero-sum game between the principal and Nature.

Our saddle-point construction yields two important corollaries. First, given the equilibrium strategy of Nature, the principal is indifferent among all contracts in the support of the optimal contract. Thus, commitment to randomization is unnecessary. Second, using a more general mechanism that asks the agent to report the technology and conditions the choice of contract on the report cannot improve the principal’s guarantee because the performance of any such mechanism is dictated by its payoff against the worst-case environment. Taken together, this implies that a fully general mechanism can do no better than a stochastic choice of a contract by the principal. Note also that because the saddle point is a mixed strategy Nash equilibrium, any standard interpretation of mixed strategies can be applied to this choice. For instance, one could consider a population interpretation or pursue purification.

Comparing the optimal random linear contract to the optimal deterministic linear contract reveals that the agent sometimes receives a larger share of output under the random contract. The probability of this happening can be substantial. Indeed, we provide an example where this is the case almost surely, leading the agent to take a (weakly) more productive action from any action set than under the optimal deterministic contract. Notably, the agent also unambiguously benefits from the randomization in the example. This provides a counterpoint to [Kambhampati’s \(2023\)](#) argument that showed that randomization allows the principal to extract more rents from the agent while keeping the worst-case productivity the same. In general, the principal’s gain from randomization may comprise both efficiency gains and improved rent extraction.

To gain further insight into the optimal contract, we consider environments with maximal uncertainty, i.e., those where the principal knows of only one action. After normalizing the expected output of the known action, the results for any such environment are summarized by a single parameter: the cost of the known action. This facilitates clean comparative statics as well as a comparison of the optimal random contract to the deterministic one. Among other things, we show that the gain from randomization may be arbitrarily large.

Regarding the issue of the simplicity of robustly optimal contracts, our results admit two interpretations. On one hand, the optimal random contract we identify is simple because the agent is only ever offered a linear contract. Moreover, the randomization belongs to a one-parameter family of distributions. On the other hand, a new layer of complexity is

introduced because it is optimal to randomize over a continuum of contracts. Either way, it is reassuring that with or without randomization, the robustly optimal solution is to give the agent a linear contract.

That randomization can hedge against Knightian uncertainty, or ambiguity, was suggested by [Raiffa \(1961\)](#) in response to [Ellsberg \(1961\)](#). As noted above, this is taken as the key axiom in the standard [Gilboa and Schmeidler \(1989\)](#) axiomatization of max-min utility. Decision theory has also axiomatized attitudes towards uncertainty where randomization only provides a partial hedge or, as a corner case, not at all. However, these require more complicated primitives than the standard Anscombe–Aumann framework. For example, [Saito \(2015\)](#) imposes axioms on preferences over menus of acts, whereas [Ke and Zhang \(2020\)](#) do so on preferences over randomizations over acts. Nevertheless, they also embed the Gilboa–Schmeidler preferences as a special case.

This work is related to the literature seeking foundations for linear and other commonly observed contracts following [Holmström and Milgrom \(1987\)](#), who concluded their fully Bayesian analysis of the optimality of linear intertemporal incentives by suggesting that the reason for the popularity of linear schemes might be their great robustness, and that the case for it could perhaps be made more effectively with a non-Bayesian model. We refer the reader to [Carroll \(2019\)](#) for a survey of this literature. Subsequent work has extended and generalized [Carroll’s \(2015\)](#) approach to other environments with moral hazard and established the robust optimality of linear contracts therein (see the works cited in the opening paragraph). Related models are studied also by [Antic \(2021\)](#), [Antic and Georgiadis \(2024\)](#), [Rosenthal \(2023, 2025\)](#), [Burkett and Rosenthal \(2024\)](#), and [Kambhampati \(2024\)](#), who find robustly optimal deterministic contracts different from linear contracts.

More recently, [Olszewski \(2025\)](#) shows that, if the family of possible technologies is restricted in the Carroll model, the optimality of non-linear contracts can be seen as typical in a topological sense. [Hanany and Klibanoff \(2025\)](#) consider the case where the agent can disclose hard evidence about an available action unknown to the principal. They show that optimal deterministic contracts are not robust to disclosure, whereas optimal random contracts are, thereby also highlighting the importance of allowing for randomization.

There is also a connection to recent developments in informationally robust mechanism design. Namely, [Brooks and Du \(2024, 2025\)](#) show that, quite generally, the worst-case information structure consistent with a given prior can be taken to be ordered, with only uni-directional local incentive compatibility constraints binding. If we interpret our random contract as the type distribution in Nature’s mechanism design problem, then our [Proposition 2](#) showing that any (finite) random contract is improvable with a random linear contract can be interpreted as saying that the worst-case information structure for Nature is ordered

(since linear contracts are ordered by their slopes); moreover, only upward adjacent incentive constraints bind under random linear contracts (see section 4). Our proof, however, is quite different, and in our setting there is no prior distribution.

The rest of this paper is organized as follows. Section 2 sets up the model. Section 3 establishes that any finitely supported contract can be weakly improved by a randomization over linear contracts with a weakly smaller support. Section 4 sketches how to identify optimal random contracts on a grid and supplies a guess for the form of the optimal random contract by taking the size of the grid large. Section 5 verifies the guess by constructing a saddle point in the zero-sum game between the principal and Nature. Section 6 considers the gain from randomization and Section 7 develops comparative statics under maximal uncertainty. Section 8 concludes with a discussion of corollaries and extensions of the main results. Appendices A–C contain proofs and derivations omitted from the main text.

## 2 Model

We consider the problem of a principal designing a contract to motivate an agent subject to moral hazard, with the principal facing non-quantifiable uncertainty over the agent's set of feasible actions.

Throughout, we write  $uv := \sum_{i=1}^d u_i v_i$  for the usual inner product of any two vectors  $v$  and  $u$  in  $\mathbb{R}^d$ . Moreover, for any topological space  $B$ , we use  $\Delta(B)$  to denote the set of all finitely supported probability measures on  $B$  and  $\bar{\Delta}(B)$  to denote the set of all Borel probability measures on  $B$ .

Let  $Y := \{y_1, \dots, y_n\} \subset \mathbb{R}_+$  be the finite set of possible output levels, labeled in increasing order so that  $0 \equiv y_1 < \dots < y_n$ . It will be convenient to view the possible output levels as a vector  $y = (y_1, \dots, y_n) \in \mathbb{R}_+^n$  and introduce the index set  $I := \{1, \dots, n\}$ .

An (unobservable) action for the agent is a pair  $(\pi, c) \in \Delta(Y) \times \mathbb{R}_+$ , where  $\pi = (\pi_1, \dots, \pi_n)$  is the output distribution and  $c$  is the associated cost to the agent. A *technology* is a nonempty compact set of actions  $A \subset \Delta(Y) \times \mathbb{R}_+$ .

The agent is protected by limited liability, requiring payments to him to be non-negative. A (*deterministic*) *contract* is thus a non-negative  $n$ -vector  $w = (w_1, \dots, w_n) \in \mathbb{R}_+^n$ , where  $w_i$  is the payment from the principal to the agent if output  $y_i$  is realized.

Both parties are assumed risk-neutral. Thus, if the agent plays the action  $(\pi, c)$  given contract  $w$ , then the principal's expected payoff is  $\pi(y - w) = \sum_{i \in I} \pi_i (y_i - w_i)$ , whereas the agent gets  $\pi w - c = \sum_{i \in I} \pi_i w_i - c$ .

The principal does not know the actual technology available to the agent. She is only aware of some technology  $A^0$ , referred to as the *known technology*, and views any technology

$A \supseteq A^0$  as possible. We assume throughout that the known technology  $A^0$  contains a surplus-generating action, i.e., there exists  $(\pi, c) \in A^0$  such that  $\pi y - c > 0$ .

Faced with this uncertainty, the principal evaluates contracts based on their guaranteed performance over all technologies that contain the known one. To state this formally, given a contract  $w$  and a technology  $A$ , denote the agent's maximum payoff by

$$U(w|A) := \max_{(\pi, c) \in A} \pi w - c. \quad (2.1)$$

In the special case of the known technology  $A^0$ , we write

$$U^0(w) := U(w|A^0). \quad (2.2)$$

We note for future reference that  $U^0(w)$  is continuous in the contract  $w$  by the theorem of the maximum. The corresponding payoff to the principal is

$$V(w|A) := \min_{(\pi, c) \in A} \pi(y - w) \quad \text{s.t.} \quad \pi w - c = U(w|A). \quad (2.3)$$

Note, in particular, that if the agent has multiple maximizers, then ties are broken against the principal. The principal's *guarantee* from the contract  $w$  is then

$$V(w) := \inf_{A \supseteq A^0} V(w|A). \quad (2.4)$$

We assume that the principal can hedge against the uncertainty over the agent's technology by randomizing over contracts. Formally, a *random contract* is a probability measure  $p$  on  $\mathbb{R}_+^n$ . Then,  $\Delta(\mathbb{R}_+^n)$  is the space of finitely supported random contracts and  $\bar{\Delta}(\mathbb{R}_+^n)$  is the space of all random contracts.

Given any random contract  $p$ , the agent observes the realized contract  $w$  before choosing an action. We can thus extend the principal's payoff in (2.3) to random contracts by setting

$$V(p|A) := \mathbb{E}_p[V(w|A)]. \quad (2.5)$$

Similarly, we extend the guarantee (2.4) to random contracts by setting

$$V(p) := \inf_{A \supseteq A^0} V(p|A). \quad (2.6)$$

The principal's *optimal guarantee* is the supremum of the guarantee (2.6) over all random contracts, or  $\sup_{p \in \bar{\Delta}(\mathbb{R}_+^n)} V(p)$ . A random contract is *optimal* if it attains this optimal guarantee. In contrast, the *optimal deterministic guarantee* is the supremum of (2.4) over all

contracts, or equivalently, the supremum of (2.6) over all degenerate random contracts (i.e., contracts whose support is a singleton). A contract  $w$  is an *optimal deterministic contract* if it attains the optimal deterministic guarantee. We note that the existence of a known surplus-generating action ensures that the optimal deterministic guarantee, and hence the optimal guarantee, is positive.

Linear contracts play a central role in the analysis. A deterministic contract  $w$  is *linear* if  $w = \alpha y$  for some slope  $\alpha \in [0, 1]$ . A random contract  $p \in \bar{\Delta}(\mathbb{R}_+^n)$  is *linear*, or a *random linear contract*, if every contract in the support of  $p$  is linear, i.e., for all  $w \in \text{supp}(p)$ , there exists a slope  $\alpha \in [0, 1]$  such that  $w = \alpha y$ . Whenever convenient, we identify each linear contract with its slope, and then identify the set of random linear contracts with  $\bar{\Delta}([0, 1])$ .

Some remarks are in order regarding the formulation of the problem:

1. Because we assume that the agent observes the realized contract before choosing an action, the agent need not know the underlying random contract, or even be aware that the contract was generated via randomization. Indeed, because of bilateral risk neutrality, delaying the realization of the random contract until after the agent has chosen an action would not be useful: Given any action  $(\pi, c)$ , the principal's and the agent's payoffs from a randomized contract  $p$  would be  $\pi(y - \mathbb{E}_p[w])$  and  $\pi \mathbb{E}_p[w] - c$ , and thus we could equivalently use the deterministic contract  $\tilde{w} := \mathbb{E}_p[w]$ .
2. Because the agent knows the true technology, it may seem restrictive to not allow the principal to use contracts that ask the agent to report this information. However, it follows as an immediate corollary of the saddle-point construction in the proof of Proposition 3 that using such contracts cannot improve the principal's optimal guarantee; see Section 8.3.
3. The principal can be seen as a decision maker in the framework of Anscombe and Aumann (1963) whose preferences satisfy the axioms of Gilboa and Schmeidler (1989). The key among them is the uncertainty-aversion axiom. Translated to the present context, it states that, when indifferent between contracts, the principal prefers any randomization between them to choosing either one deterministically. We develop this connection formally in Appendix C.
4. Our contracting environment is the same as in Carroll (2015), save for the following minor differences. First, we take the set of outputs,  $Y$ , to be finite (rather than just compact) to minimize technicalities. This is not required to prove Proposition 3, but it allows for simple duality-based proofs in Section 3. Second, we assume adversarial tie-breaking in (2.3), whereas Carroll broke ties in the principal's favor. Adversarial

tie-breaking could be argued to better capture the spirit of the robustness exercise, and it simplifies the analysis as the infimum in (2.6) becomes a minimum. However, as we will see, both assumptions lead to the same optimal guarantee and the optimal contract is similarly unaffected by tie-breaking except for an uninteresting corner case.

### 3 Linear improvement argument

The duality approach we adopt to prove the optimality of random linear contracts also yields a short proof of the optimality of linear contracts within the class of deterministic contracts. It is instructive to see the argument first in this simpler case. We then generalize from the deterministic case by showing that, for any finite random contract, there exists a random linear contract with weakly smaller support that obtains at least the same payoff guarantee.

#### 3.1 Deterministic case

Consider a deterministic contract  $w$ . We formulate the guarantee (2.4) of  $w$  as a linear program. To this end, note that, given any technology  $A \supseteq A^0$ , only the action chosen by the agent matters for the principal's payoff  $V(w|A)$  in (2.3).<sup>2</sup> It thus suffices to take the infimum in (2.4) over technologies that add at most one new action to the known technology  $A^0$ . It follows that the contract's guarantee is characterized by the following linear program:

$$V(w) = \min_{\pi, c} \pi(y - w) \tag{3.1}$$

$$\text{s.t. } \pi w - c \geq U^0(w), \tag{3.2}$$

$$\sum_{i=1}^n \pi_i = 1, \tag{3.3}$$

$$c, \pi_i \geq 0 \quad \forall i \in I. \tag{3.4}$$

That is, Nature designs an action  $(\pi, c)$  to minimize the principal's profit, with constraint (3.2) ensuring that  $(\pi, c)$  is a best response for the agent since, by (2.2),  $U^0(w)$  is the agent's maximum payoff from the known technology  $A^0$ . Problem (3.1–3.4) is feasible because taking  $(\pi, c)$  to be an action in  $A^0$  that attains  $U^0(w)$  satisfies all constraints.

We can now show that any contract can be weakly improved upon by some linear contract.

**Proposition 1.** *For every contract  $w$ , there is a linear contract  $\alpha$  such that  $V(w) \leq V(\alpha)$ .*

---

<sup>2</sup>I.e., if  $(\pi^*, c^*)$  attains the minimum in (2.3) given  $w$  and  $A \supseteq A^0$ , then  $V(w|A) = V(w|A^0 \cup \{(\pi^*, c^*)\})$ .

*Proof.* Fix a contract  $w$ . The guarantee  $V(w)$  is clearly bounded from below by  $-\max_i w_i$ . By strong duality, the following dual to problem (3.1–3.4) is then feasible and bounded:

$$V(w) = \max_{\lambda, \mu} \lambda U^0(w) + \mu \quad (3.5)$$

$$\text{s.t. } \lambda w_i + \mu \leq y_i - w_i \quad \forall i \in I, \quad (3.6)$$

$$\lambda \geq 0. \quad (3.7)$$

Let  $(\lambda^*, \mu^*)$  be an optimal solution. Evaluating (3.6) at  $i = 1$  gives  $\mu^* \leq -(1 + \lambda^*)w_1 \leq 0$  as  $y_1 = 0$ . Define a new contract  $\hat{w}$  as follows. For each  $i$ , choose  $\hat{w}_i$  to satisfy (3.6) as an equality, i.e., let

$$\hat{w}_i := \frac{y_i - \mu^*}{\lambda^* + 1}. \quad (3.8)$$

By inspection,  $\hat{w}$  is an affine function of output. As the original contract  $w$  also satisfies (3.6), we have  $\hat{w} \geq w \geq 0$ . Thus,  $\hat{w}$  is a well-defined contract. Moreover, we have

$$V(\hat{w}) \geq \lambda^* U^0(\hat{w}) + \mu^* \geq \lambda^* U^0(w) + \mu^* = V(w), \quad (3.9)$$

where the first inequality follows because  $(\lambda^*, \mu^*)$  is by construction of  $\hat{w}$  still feasible when  $w$  is replaced with  $\hat{w}$  in the dual, and the second inequality follows because  $\hat{w} \geq w$  and  $U^0(w)$  is weakly increasing in  $w$  by inspection of (2.2). We conclude that the original contract  $w$  is outperformed by the affine contract  $\hat{w}$ .

To get a linear contract, we may drop the lump-sum payment  $-\mu^*/(\lambda^* + 1) \geq 0$  from  $\hat{w}$  by letting  $\hat{\hat{w}} := y/(\lambda^* + 1)$ . This leaves the agent's choice from every technology unchanged, but weakly increases the principal's payoff. Thus,  $V(\hat{\hat{w}}) \geq V(\hat{w}) \geq V(w)$ .  $\square$

A reader familiar with Carroll's (2015) proof will see the parallels between the arguments, with duality here replacing the separating hyperplane theorem. Specifically, the affine contract  $\hat{w}$  defined by (3.8) is a supporting hyperplane to the convex hull of the graph  $\{(y_i, w_i) : i \in I\}$  of the original contract  $w$ . Given the equivalence of linear programming duality and the separating hyperplane theorem, at a deeper level the proofs are the same. However, the linear programming perspective allows for a somewhat more concise argument.<sup>3</sup>

Proposition 1 implies that linear contracts are optimal, in the following sense:

**Corollary 1.** *The optimal deterministic guarantee satisfies  $\sup_{w \in \mathbb{R}_+^n} V(w) = \sup_{\alpha \in [0,1]} V(\alpha)$ .*

<sup>3</sup>The proof of Proposition 1 extends mutatis mutandis to any compact set  $Y \subset \mathbb{R}_+$  such that  $\min Y = 0$ , with a contract then defined to be a continuous function  $w : Y \rightarrow \mathbb{R}_+$ . The primal problem (3.1–3.4) then becomes one of choosing a regular nonnegative Borel measure  $\pi$  on  $Y$  and a cost  $c \in \mathbb{R}$  to solve

$$V(w) = \min \int_Y (y - w(y)) d\pi(y) \quad \text{s.t.} \quad \int_Y w(y) d\pi(y) - c \geq U^0(w), \quad \int_Y 1 d\pi(y) = 1, \quad \text{and } c \geq 0.$$

*Proof.* Clearly  $\bar{v} := \sup_{w \in \mathbb{R}_+^n} V(w) \geq \sup_{\alpha \in [0,1]} V(\alpha)$ . For the converse, let  $w_n$  be a sequence such that  $V(w_n) \rightarrow \bar{v}$ . By Proposition 1, there is a sequence of linear contracts  $\alpha_n$  such that  $V(w_n) \leq V(\alpha_n) \leq \sup_{m \geq n} V(\alpha_m)$  for all  $n$ . Hence,  $\bar{v} \leq \limsup_n V(\alpha_n) \leq \sup_{\alpha \in [0,1]} V(\alpha)$ .  $\square$

Having established the optimality of linear contracts as a class, it only remains to characterize the optimal deterministic guarantee and the optimal linear contracts to conclude the analysis of deterministic contracts. This task is accomplished in Appendix A, which largely follows from the analysis in Carroll (2015), but serves to illustrate that adversarial and principal-optimal tie-breaking yield essentially the same results. Specifically, the guarantee of any non-zero linear contract is the same under both tie-breaking rules. The only difference is that under principal-optimal tie-breaking, the zero contract may have a positive guarantee, whereas with adversarial tie-breaking its guarantee is zero. However, the guarantee from the principal-optimal case can still be obtained in the limit of any sequence of linear contracts with slopes converging to zero.

### 3.2 Random contracts

We now consider random contracts. To facilitate our duality-based arguments, we temporarily restrict attention to finitely supported random contracts (which are dense in the space of all contracts). The following results generalize Proposition 1 and Corollary 1 to such contracts.

**Proposition 2.** *For every finitely supported random contract  $p \in \Delta(\mathbb{R}_+^n)$ , there is a finitely supported random linear contract  $q \in \Delta(\mathbb{R}_+^n)$  such that*

$$V(p) \leq V(q) \text{ and } |\text{supp}(p)| \geq |\text{supp}(q)|.$$

**Corollary 2.** *The optimal guarantee from finitely supported random contracts satisfies*

$$\sup_{p \in \Delta(\mathbb{R}_+^n)} V(p) = \sup_{p \in \Delta([0,1])} V(p).$$

*Proof.* Same as Corollary 1, mutatis mutandis.  $\square$

The dual problem (3.5–3.7) in turn becomes one of choosing numbers  $\lambda \in \mathbb{R}$  and  $\mu \in \mathbb{R}$  to

$$\max \lambda U^0(w) + \mu \text{ s.t. } \lambda w(y) + \mu \leq y - w(y) \ \forall y \in Y, \text{ and } \lambda \geq 0.$$

As the zero contract has  $V(0) \geq 0$ , we may assume that  $V(w) > 0$ . By weak duality, this implies that the dual is also bounded. Moreover, if we let  $\lambda' = 1$  and  $\mu' = \min\{y - 2w(y) : y \in Y\} - 1$ , then  $(\lambda', \mu')$  satisfies all dual constraints as strict inequalities. Therefore, strong duality holds for this pair of semi-infinite linear programs (see, e.g., Lai and Wu, 1992, Theorem 2.1). The rest of the argument now proceeds as above, with the affine improvement contract constructed analogously to (3.8) given some optimal dual solution.

We note that because of the conclusion regarding supports, Proposition 2 nests Proposition 1 as a special case by taking  $p$  to be a degenerate random contract. More generally, Proposition 2 implies that random linear contracts are optimal given any upper bound on the size of a contract's support (e.g., because of computational or complexity considerations).

In order to establish Proposition 2, the first step is to formulate the guarantee (2.6) as a linear program. Let  $p \in \Delta(\mathbb{R}_+^n)$  be a finitely supported random contract. Enumerate the contracts in the support of  $p$  so that  $\text{supp}(p) = \{w^1, \dots, w^k\}$  and denote the index set by  $T := \{1, \dots, k\}$ . We will write  $p_t := p(w^t)$  for the probability of contract  $t$ . In searching for a worst-case technology against  $p$ , it suffices to consider technologies  $A \supseteq A^0$  that have at most  $k$  new actions, one for each possible realized contract.<sup>4</sup> The guarantee  $V(p)$  is thus characterized by the following finite linear program, which we refer to as the primal problem:

$$V(p) = \min_{\{(\pi^t, c_t)\}} \sum_{t \in T} p_t \pi^t(y - w^t) \quad (3.10)$$

$$\text{s.t.} \quad \pi^t w^t - c_t \geq \pi^s w^t - c_s \quad \forall t, s \in T : t \neq s, \quad (\kappa_{ts}) \quad (3.11)$$

$$\pi^t w^t - c_t \geq U^0(w^t) \quad \forall t \in T, \quad (\lambda_t) \quad (3.12)$$

$$\sum_{i \in I} \pi_i^t = 1 \quad \forall t \in T, \quad (\mu_t) \quad (3.13)$$

$$c_t, \pi_i^t \geq 0 \quad \forall i \in I, t \in T. \quad (3.14)$$

That is, Nature designs, for each contract  $w^t$  in the support of  $p$ , an action  $(\pi^t, c_t)$  the agent will take if contract  $w^t$  is realized. This problem can be interpreted as a multi-dimensional mechanism design problem where the agent's type  $t$  corresponds to a contract  $w^t \in \mathbb{R}_+^n$ , the allocation is an output distribution  $\pi \in \Delta(Y)$ , and the transfer is a cost  $c$ , constrained to be nonnegative. The constraint (3.12) is a participation constraint that ensures that each type  $t$  prefers its own action  $(\pi^t, c_t)$  to any of the known actions in  $A^0$ . (Note that the utility from this outside option depends on the type  $t$ .) The new constraint (3.11), absent from Nature's problem studied in the deterministic case, is an incentive compatibility constraint that ensures that each type  $t$  prefers its own action  $(\pi^t, c_t)$  to the action  $(\pi^s, c_s)$  of every other type  $s$ .

Problem (3.10–3.14) is feasible because taking each  $(\pi^t, c_t)$  to be an action in  $A^0$  that

---

<sup>4</sup>To see this, given any technology  $A \supseteq A^0$ , let  $a^t := (\pi^t, c_t)$  denote the action the agent chooses from  $A$  if contract  $w^t$  is realized. (That is,  $a^t$  attains  $V(w^t|A)$ .) It is straightforward to verify that, for all  $t \in T$ , we have  $V(w^t|A) = V(w^t|A^0 \cup \{a^1, \dots, a^k\})$ . This implies that

$$V(p|A) = \sum_{t \in T} p_t V(w^t|A) = \sum_{t \in T} p_t V(w^t|A^0 \cup \{a^1, \dots, a^k\}) = V(p|A^0 \cup \{a^1, \dots, a^k\}).$$

Therefore, for any  $A \supseteq A^0$ , there exists a technology  $A' \supseteq A^0$  with  $|A' \setminus A^0| \leq k$  such that  $V(p|A') = V(p|A)$ .

attains  $U^0(w^t)$  satisfies all constraints. (Put differently, Nature can always just give the agent the known technology  $A^0$ .) Moreover, its value is clearly bounded from below by  $-\max_{i,t} w_i^t$ , and thus an optimal solution exists.

As in the deterministic case, we will study the dual to Nature's problem. Strong duality implies that the following dual is feasible and bounded, with  $V(p)$  the optimal value:

$$V(p) = \max_{\kappa, \lambda, \mu} \sum_{t \in T} (\lambda_t U^0(w^t) + \mu_t) \quad (3.15)$$

$$\text{s.t. } \lambda_t w_i^t + \mu_t + \sum_{s \neq t} \kappa_{ts} w_i^t - \sum_{s \neq t} \kappa_{st} w_i^s \leq p_t(y_i - w_i^t) \quad \forall i \in I, t \in T, \quad (\pi_i^t) \quad (3.16)$$

$$\lambda_t + \sum_{s \neq t} \kappa_{ts} - \sum_{s \neq t} \kappa_{st} \geq 0 \quad \forall t \in T, \quad (c_t) \quad (3.17)$$

$$\kappa_{ts}, \lambda_t \geq 0 \quad \forall t, s \in T : t \neq s. \quad (3.18)$$

We will use *dual solution* to refer to any vector  $(\kappa, \lambda, \mu) \in \mathbb{R}^{T(T-1)} \times \mathbb{R}^T \times \mathbb{R}^T$  satisfying constraints (3.17) and (3.18), and use the qualifiers *feasible* or *optimal* to indicate, respectively, that the dual solution is feasible or optimal in (3.15–3.18) for a given random contract  $p \in \Delta(\mathbb{R}_+^n)$ .

In order to make duality between problems (3.10–3.14) and (3.15–3.18) easier to verify, each non-trivial primal constraint in (3.10–3.14) has the associated dual variable displayed next to it in parenthesis. Similarly, the associated primal variable is displayed next to each non-trivial dual constraint in (3.15–3.18). It may also be instructive to compare the above dual problem to the dual (3.5–3.7) from the deterministic case; by inspection, the problems coincide if the random contract  $p$  is degenerate so that  $T$  is a singleton.

The general idea in the proof of Proposition 2 is to proceed analogously to the proof of Proposition 1 from the deterministic case and use constraint (3.16) to identify affine improvement contracts, which can then be linearized to obtain a further improvement. The complication in the random case is that the  $(i, t)$ -instance of constraint (3.16) features not just contract  $w^t$  itself, but also every contract  $w^s$  for which the  $(s, t)$ -instance of the incentive compatibility constraint (3.11) is binding. To handle this issue, we first show that an improvement argument analogous to the deterministic case goes through for any contract  $w^t$  for which these other contracts  $w^s$  are linear. We then prove a structural result about the acyclicity of optimal dual solutions that allows this argument to be iterated to linearize all contracts in the support of  $p$ .

It is useful to introduce the following definitions to describe the relevant structure of optimal dual solutions and binding incentive constraints. Given a dual solution  $(\kappa, \lambda, \mu)$ , let  $G(\kappa)$  be a directed graph with vertex set  $T = \{1, \dots, k\}$  and an arc directed from  $t$  to

$s$  whenever  $\kappa_{ts} > 0$ . Given a vertex  $t \in T$ , we refer to  $N(t) := \{s \in T : \kappa_{st} > 0\}$  as the set of *in-neighbors* of  $t$ . Similarly,  $O(t) := \{s \in T : \kappa_{ts} > 0\}$  is the set of *out-neighbors* of  $t$ . Note that if  $(\kappa, \lambda, \mu)$  is an optimal solution to (3.15–3.18), then by complementary slackness, each arc in  $G(\kappa)$  corresponds to a binding instance of the incentive compatibility constraint (3.11) in the primal problem. Specifically,  $s$  is an in-neighbor of  $t$  if the incentive compatibility constraint keeping type  $s$  from mimicking type  $t$  binds (i.e., if keeping the agent with contract  $w^s$  from choosing the action  $(\pi^t, c_t)$  intended for contract  $w^t$  is a binding constraint), whereas  $s$  is an out-neighbor of  $t$  if the constraint keeping  $t$  from mimicking  $s$  binds.

The following lemma forms the heart of the proof of Proposition 2. Its proof closely parallels that of Proposition 1 from the deterministic case.

**Lemma 1.** *Let  $p$  be a random contract with finite support  $\{w^t : t \in T\}$ , and let  $(\kappa^*, \lambda^*, \mu^*)$  be an optimal dual solution. Suppose all in-neighbors of a contract  $\tau \in T$  in the graph  $G(\kappa^*)$  are linear, i.e., the contract  $w^s$  is linear for all  $s \in N(\tau)$ . Then there exists a linear contract  $\alpha_\tau$  and a random contract  $q$  with support  $\{w^t : t \in T \setminus \{\tau\}\} \cup \{\alpha_\tau\}$  such that  $V(q) \geq V(p)$ .*

*Proof.* Fix a finitely supported contract  $p$  and an optimal dual solution  $(\kappa^*, \lambda^*, \mu^*)$ . Let  $\tau \in T$  be the index of a contract whose every in-neighbor  $s \in N(\tau)$  is linear with some slope  $\alpha_s \in [0, 1]$ . The  $(i, \tau)$ -instance of constraint (3.16) then takes the form

$$\lambda_\tau^* w_i^\tau + \mu_\tau^* + \sum_{s \in O(\tau)} \kappa_{\tau s}^* w_i^\tau - \sum_{s \in N(\tau)} \kappa_{s\tau}^* \alpha_s y_i \leq p_\tau (y_i - w_i^\tau),$$

or, equivalently,

$$w_i^\tau \leq \frac{p_\tau + \sum_{s \in N(\tau)} \kappa_{s\tau}^* \alpha_s}{\lambda_\tau^* + p_\tau + \sum_{s \in O(\tau)} \kappa_{\tau s}^*} y_i - \frac{\mu_\tau^*}{\lambda_\tau^* + p_\tau + \sum_{s \in O(\tau)} \kappa_{\tau s}^*} =: \hat{w}_i^\tau. \quad (3.19)$$

As  $i$  ranges over  $I$ , the right-hand side defines a vector  $\hat{w}^\tau$ . By construction,  $\hat{w}^\tau \geq w^\tau \geq 0$ , and thus  $\hat{w}^\tau$  is a contract. By inspection, it is affine in output, i.e., of the form  $\hat{w}_i^\tau = \alpha_\tau y_i + \beta_\tau$ . The constant  $\beta_\tau$  is nonnegative, because  $y_1 = 0$  implies that  $\beta_\tau = \hat{w}_1^\tau \geq w_1^\tau \geq 0$ . To see that the slope  $\alpha_\tau$  is in  $[0, 1]$ , we note that

$$0 < \frac{p_\tau + \sum_{s \in N(\tau)} \kappa_{s\tau}^* \alpha_s}{\lambda_\tau^* + p_\tau + \sum_{s \in O(\tau)} \kappa_{\tau s}^*} \leq \frac{p_\tau + \sum_{s \in N(\tau)} \kappa_{s\tau}^*}{\lambda_\tau^* + p_\tau + \sum_{s \in O(\tau)} \kappa_{\tau s}^*} \leq \frac{p_\tau}{p_\tau} = 1,$$

where the first inequality follows because  $p_\tau > 0$  and  $\alpha_s, \kappa_{s\tau}^*, \kappa_{\tau s}^*, \lambda_\tau^* \geq 0$ , the second follows because  $\alpha_s \leq 1$ , and the third inequality follows because  $(\kappa^*, \lambda^*, \mu^*)$  satisfies (3.17).

Let  $V(p|w^\tau \leftarrow \hat{w}^\tau)$  denote the value of the dual problem (3.15–3.18) for contract  $p$  when

the contract  $w^\tau$  is replaced with the affine contract  $\hat{w}^\tau$  just constructed. We observe that the optimal dual solution  $(\kappa^*, \lambda^*, \mu^*)$  remains feasible in this problem. To see this, note that it satisfies the  $(i, \tau)$ -instance of (3.16) with equality for all  $i \in I$  by construction of  $\hat{w}^\tau$ . Moreover, because  $\hat{w}^\tau \geq w^\tau$  and  $\kappa_{s\tau}^* \geq 0$  for all  $s \neq \tau$ , replacing  $w^\tau$  with  $\hat{w}^\tau$  weakly relaxes all other instances of (3.16). Therefore, we have

$$V(p|w^\tau \leftarrow \hat{w}^\tau) \geq \sum_{t \in T \setminus \{\tau\}} \lambda_t^* U^0(w^t) + \lambda_\tau^* U^0(\hat{w}^\tau) + \sum_{t \in T} \mu_t^* \geq \sum_{t \in T} (\lambda_t^* U^0(w^t) + \mu_t^*) = V(p), \quad (3.20)$$

where the first inequality is by feasibility of  $(\kappa^*, \lambda^*, \mu^*)$ , the second inequality follows because  $\hat{w}^\tau \geq w^\tau$  by construction and because  $U^0$  is weakly increasing in the contract (in the pointwise order) by inspection of (2.2), and the equality is by optimality of  $(\kappa^*, \lambda^*, \mu^*)$  for the original random contract  $p$ . We conclude that replacing  $w^\tau$  with  $\hat{w}^\tau$  yields at least a weak improvement in the dual value.

To complete the proof, note that by strong duality,  $V(p|w^\tau \leftarrow \hat{w}^\tau)$  is the optimal value of the primal problem (3.10–3.14), which now takes the form

$$\begin{aligned} V(p|w^\tau \leftarrow \hat{w}^\tau) = \min_{\{(\pi^t, c_t)\}} & p_\tau(1 - \alpha_\tau)\pi^\tau y - p_\tau\beta_\tau + \sum_{t \in T \setminus \{\tau\}} p_t\pi^t(y - w^t) \\ \text{s.t.} & \pi^t w^t - c_t \geq \pi^s w^t - c_s & \forall t, s \in T : t \neq s, \tau \\ & \pi^t w^t - c_t \geq U^0(w^t) & \forall t \in T : t \neq \tau, \\ & \alpha_\tau \pi^\tau y + \beta_\tau - c_\tau \geq \alpha_\tau \pi^s y + \beta_\tau - c_s & \forall s \in T : s \neq \tau, \\ & \alpha_\tau \pi^\tau y + \beta_\tau - c_\tau \geq U^0(\alpha_\tau) + \beta_\tau \\ & \sum_{i \in I} \pi_i^t = 1 & \forall t \in T, \\ & c_t, \pi_i^t \geq 0 & \forall i \in I, t \in T. \end{aligned} \quad (3.21)$$

Note that the feasible set is independent of the constant  $\beta_\tau \geq 0$ . (It enters both sides of the incentive compatibility and participation constraints for type  $\tau$  and cancels out.) Thus, we may further improve the value by replacing the affine contract  $\hat{w}^\tau$  with the linear contract  $\alpha_\tau$ . Specifically, letting  $V(p|w^\tau \leftarrow \alpha_\tau)$  denote the optimal value for the above primal program (3.21) when  $\beta_\tau \equiv 0$ , we have  $V(p|w^\tau \leftarrow \alpha_\tau) = V(p|w^\tau \leftarrow \hat{w}^\tau) + p_\tau\beta_\tau \geq V(p)$ , where the inequality is by (3.20).

Only some housekeeping remains: If the linear contract  $\alpha_\tau$  is different from all the contracts  $w^t$ ,  $t \in T \setminus \{\tau\}$ , then it is immediate that  $V(p|w^\tau \leftarrow \alpha_\tau)$  is the value of a random contract with support  $\{w^t : t \in T \setminus \{\tau\}\} \cup \{\alpha_\tau\}$  that assigns probability  $p_t$  to each contract  $w^t$  with  $t \neq \tau$  and probability  $p_\tau$  to contract  $\alpha_\tau$ . If instead  $\alpha_\tau = w^{\hat{t}}$  for some  $\hat{t} \neq \tau$ , then

we can define a random contract  $q$  by setting  $q(w^{\hat{t}}) := p(w^{\hat{t}}) + p(w^\tau)$  and  $q(w^t) := p(w^t)$  for all  $t \neq \hat{t}, \tau$ . Then  $\text{supp}(q) = \{w^t : t \in T \setminus \{\tau\}\} = \{w^t : t \in T \setminus \{\tau\}\} \cup \{\alpha_\tau\}$  as desired. Furthermore,  $V(q)$  is the optimal value of the problem (3.21) with  $\beta_\tau \equiv 0$  and with the additional constraint that  $(\pi^{\hat{t}}, c_{\hat{t}}) = (\pi^\tau, c_\tau)$ . Therefore,  $V(q) \geq V(p|w^\tau \leftarrow \alpha_\tau) \geq V(p)$ .  $\square$

Analogously to the deterministic case, the affine contract  $\hat{w}^\tau$  defined by (3.19) is a supporting hyperplane to the graph of the original contract  $w^\tau$ . To see this, note that  $w_i^\tau \geq \hat{w}_i^\tau$  for each output  $i$  by construction. Moreover, the primal variable associated with the  $(i, \tau)$ -instance of (3.16) is  $\pi_i^\tau$ , the probability of output  $i$  under action  $(\pi^\tau, c_\tau)$ . Thus, we necessarily have  $\pi_j^\tau > 0$  for some output  $j$ , and hence  $\hat{w}_j^\tau = w_j^\tau$  by complementary slackness.

In order to apply the previous lemma, we need the following result about the structure of optimal dual solutions.

**Lemma 2.** *For every finitely supported random contract  $p$ , there is an optimal dual solution  $(\kappa^*, \lambda^*, \mu^*)$  for which the graph  $G(\kappa^*)$  is acyclic.*

We prove Lemma 2 in the Appendix by first relaxing the incentive compatibility constraints (3.11) in the primal problem by subtracting an  $\varepsilon > 0$  from the right-hand side and showing that every optimal dual solution to this perturbed problem is acyclic. To see the basic intuition, suppose that the  $\varepsilon$ -relaxation of the incentive compatibility constraint (3.11) holds with equality in both directions between two equiprobable types  $s$  and  $t$  under some optimal solution. Then each of these types strictly prefers the other type's bundle to its own. If we swap their bundles, then all constraints remain satisfied as we have not introduced any new deviations. Moreover, this new solution attains a strictly lower value of the objective (3.10) as total surplus is unchanged (since  $s$  and  $t$  are equiprobable), but more of it now goes to the agent as types  $s$  and  $t$  are strictly better off. Thus, the original solution is not optimal, a contradiction. The proof generalizes this logic to cycles involving any number of types that need not be equiprobable. The argument is then completed by using upper hemi-continuity of the dual solutions to obtain an acyclic solution to the unperturbed problem.

Lemma 2 allows us to identify some contract in the support of  $p$  which can be linearized to obtain an improvement using Lemma 1:

**Lemma 3.** *Let  $p$  be a finitely supported random contract whose support contains at least one non-linear contract. Then there exists a random contract  $q$  such that (i) the support of  $q$  contains one fewer non-linear contract than the support of  $p$ , (ii)  $|\text{supp}(q)| \leq |\text{supp}(p)|$ , and (iii)  $V(q) \geq V(p)$ .*

*Proof.* Let  $p$  be finitely supported and assume its support  $\{w^t : t \in T\}$  contains at least one non-linear contract. By Lemma 2, we can fix an optimal dual solution  $(\kappa^*, \lambda^*, \mu^*)$  for  $p$  such that  $G(\kappa^*)$  is acyclic. We then partition  $T$  recursively as follows:

- Let  $L_0 := \{t \in T : N(t) = \emptyset\}$ .
- For  $\ell \geq 1$ , let  $L_\ell := \{t \in T \setminus \cup_{m < \ell} L_m : N(t) \subseteq \cup_{m < \ell} L_m\}$ .

That is, the zero layer  $L_0$  consists of all contracts that have no in-neighbors in the graph  $G(\kappa^*)$ . For  $\ell > 0$ , the  $\ell$ -th layer  $L_\ell$  consists of all contracts not contained in any of the lower layers  $L_m$ ,  $m < \ell$ , but whose in-neighbors are in these layers. Because  $G(\kappa^*)$  is acyclic, it is straightforward to verify that there exists  $\bar{\ell} \in \mathbb{N}_0$  such that  $L_\ell$  is nonempty if and only if  $0 \leq \ell \leq \bar{\ell}$ , and that  $\{L_0, \dots, L_{\bar{\ell}}\}$  is a partition of  $T$ .<sup>5</sup>

Let  $L_n$  be the lowest layer that contains some non-linear contract  $w^\tau$ . Such a layer exists, because the support of  $p$  contains a non-linear contract by assumption. By definition of  $L_n$ , all in-neighbors of the non-linear contract  $w^\tau$  are in layers  $L_\ell$  for  $\ell < n$  and they are thus linear by definition of  $L_n$  (vacuously so if  $n = 0$ ).

Applying Lemma 1 to the non-linear contract  $w^\tau$  with linear in-neighbors gives us a linear contract  $\alpha_\tau$  and a random contract  $q$  with support  $\{w^t : t \in T \setminus \{\tau\}\} \cup \{\alpha_\tau\}$  such that  $V(q) \geq V(p)$ . As the support of  $q$  omits  $w^\tau$ , it contains one fewer non-linear contract than the support of  $p$ . Moreover, we have  $|\text{supp}(q)| \leq |\text{supp}(p)|$  by construction.  $\square$

We are now ready to prove Proposition 2.

*Proof of Proposition 2.* Let  $p$  be a finitely supported random contract. If  $p$  is linear, then the claim holds vacuously with  $q = p$ . So suppose the support of  $p$  contains  $m$  non-linear contracts for  $m \geq 1$ . Repeated application of Lemma 3 then yields contracts  $q^1, \dots, q^m$  such that  $V(p) \leq V(q^1) \leq \dots \leq V(q^m)$ ,  $|\text{supp}(p)| \geq |\text{supp}(q^1)| \geq \dots \geq |\text{supp}(q^m)|$ , and the support of each  $q^n$  contains  $m - n$  non-linear contracts. In particular,  $q^m$  is linear, so the claim follows by taking  $q = q^m$ .  $\square$

The proof of Proposition 2 shows that any random contract  $p \in \Delta(\mathbb{R}_+^n)$  can be (weakly) improved by linearizing each contract in its support while keeping the probabilities fixed. The argument thus reveals that the problem of hedging against Nature can be decoupled from the problem of choosing the form of the contracts given to the agent.

The new challenge in the random case is the presence of the incentive compatibility constraints (3.11). They create interlinkages between contracts because the action  $(\pi^t, c_t)$

---

<sup>5</sup>Because  $G_0 := G(\kappa^*)$  is acyclic, there exists a contract without in-neighbors. Thus,  $L_0$  is nonempty. Now construct a graph  $G_1$  by removing all contracts in  $L_0$  from the graph  $G_0$ . Since we just removed vertices,  $G_1$  is acyclic and thus contains a contract without in-neighbors. Moreover, as any such contract was not removed in the previous step, it must have an in-neighbor in  $L_0$ . Thus,  $L_1$  is nonempty. We may now construct a graph  $G_2$  by removing all contracts in  $L_1$  from  $G_1$ , and then identify  $L_2$  as the (nonempty) set of contracts without in-neighbors in  $G_2$ . Continuing iteratively, this process terminates in finitely many steps as the original graph  $G_0$  is finite.

Nature targets to contract  $w^t$  will be available to the agent also when any other contract is realized. Because the primal (3.10–3.14) is a multi-dimensional screening problem, the structure of binding incentive compatibility constraints (i.e., those with a positive shadow price) is not clear a priori. Our proof deals with this by showing that there always exists an optimal solution under which the binding constraints do not form a cycle (Lemma 2). This is enough to identify some non-linear contract in the support that can be linearized (Lemma 3); the proof is then completed by iterating this argument. A subtlety worth noting is that, at each iteration, the optimal dual solution  $(\kappa^*, \lambda^*, \mu^*)$  and hence the graph  $G(\kappa^*)$  may well be different. Nevertheless, as long as the random contract is non-linear, the graph contains some non-linear contract with only linear in-neighbors. Thus, Lemma 1 applies.

## 4 Identifying optimal contracts

We now sketch how to identify an optimal randomization over a grid of linear contracts. As the grid becomes dense in the unit interval, the limit of the corresponding optimal randomizations is a natural guess for the optimal random contract. This guess is verified to be optimal in Proposition 3 of Section 5. Because the proof of Proposition 3 does not rely on any results in this section, we proceed here somewhat informally.<sup>6</sup>

Let us specialize Nature’s primal problem (3.10–3.14) to the case of a random linear contract  $p$ . It will be convenient to let the support of  $p$  be a subset of an equally-spaced grid  $\{\alpha_1, \dots, \alpha_k\} \subset [0, 1]$ , where the possible slopes  $0 < \alpha_1 < \dots < \alpha_k \equiv 1$  are defined by  $\alpha_t := t/k$  for  $t = 0, \dots, k$ . (Thus,  $\alpha_{t+1} - \alpha_t = k^{-1}$ .) By inspection, the problem depends on each output distribution  $\pi^t$  only via the expected output  $e_t := \pi^t y \in [0, y_n]$ . With this change of variables, we can write the primal problem as follows:

$$V(p) = \min_{\{(e_t, c_t)\}} \sum_{t \in T} p_t (1 - \alpha_t) e_t \quad (4.1)$$

$$\text{s.t.} \quad \alpha_t e_t - c_t \geq \alpha_t e_s - c_s \quad \forall t, s \in T : t \neq s, \quad (4.2)$$

$$\alpha_t e_t - c_t \geq U^0(\alpha_t) \quad \forall t \in T, \quad (4.3)$$

$$e_t \leq y_n \quad \forall t \in T, \quad (4.4)$$

$$e_t, c_t \geq 0 \quad \forall t \in T. \quad (4.5)$$

Nature’s multi-dimensional mechanism design problem has thus been reduced to a one-dimensional problem: the agent’s type is a slope  $\alpha$ , the allocation is an expected output  $e$ , the “transfer” is a cost  $c$ , and the outside option is still type-dependent.

---

<sup>6</sup>See the working paper [Kambhampati, Toikka, and Vohra \(2024\)](#) for further details.

Standard arguments allow us to simplify the constraints (4.2–4.5). Specifically, the upward adjacent constraints in (4.2) can be taken to hold with equality without loss of optimality. This allows us to recursively eliminate the costs up to  $c_1$ , which can then be set to zero without loss of optimality. Adding in the usual monotonicity constraint on the allocation rule to ensure incentive compatibility and dropping the upper-bound constraints (without loss of optimality) yields a simplified program:

$$V(p) = \min_{\{e_t\}} \sum_{t \in T} p_t(1 - \alpha_t)e_t \quad (4.6)$$

$$\text{s.t.} \quad k^{-1} \sum_{s=1}^t e_s \geq U^0(\alpha_t) \quad \forall t \in T, \quad (\lambda_t) \quad (4.7)$$

$$e_{t+1} - e_t \geq 0 \quad \forall t \in T \setminus \{k\}, \quad (\theta_t) \quad (4.8)$$

$$e_t \geq 0 \quad \forall t \in T. \quad (4.9)$$

The principal chooses  $p$  to maximize  $V(p)$ . This max-min problem can be transformed into a maximization problem by replacing Nature's minimization problem (4.6–4.9) with its dual. The resulting problem characterizing the optimal guarantee over the grid is

$$\max_{p, \lambda, \theta} \sum_{t \in T} \lambda_t U^0(\alpha_t) \quad (4.10)$$

$$\text{s.t.} \quad k^{-1} \sum_{s=t}^k \lambda_s + \theta_{t-1} - \theta_t \leq p_t(1 - \alpha_t) \quad \forall t \in T, \quad (e_t) \quad (4.11)$$

$$\sum_{t \in T} p_t = 1 \quad (\delta) \quad (4.12)$$

$$\lambda_t, p_t \geq 0 \quad \forall t \in T, \quad (4.13)$$

$$\theta_t \geq 0 \quad \forall t \in T \setminus \{k\}, \quad (4.14)$$

where  $\alpha_0 \equiv 0$ ,  $\theta_0 \equiv 0$ , and  $\theta_k \equiv 0$ . By strong duality, the value of (4.10–4.14), i.e., the optimal guarantee over the grid, is simply

$$\min_{\{e_t\}, \delta} \delta \quad (4.15)$$

$$\text{s.t.} \quad \delta - (1 - \alpha_t)e_t \geq 0 \quad \forall t \in T, \quad (p_t) \quad (4.16)$$

$$k^{-1} \sum_{s=1}^t e_s \geq U^0(\alpha_t) \quad \forall t \in T, \quad (\lambda_t) \quad (4.17)$$

$$e_{t+1} - e_t \geq 0 \quad \forall t \in T \setminus \{k\}, \quad (\theta_t) \quad (4.18)$$

$$e_t \geq 0 \quad \forall t \in T. \quad (4.19)$$

Identifying a solution to (4.15–4.19) is straightforward. Note first that (4.16) can be taken to hold with equality for all  $t < k$  without loss of optimality; increasing each  $e_t$  so that the constraint binds relaxes the participation constraints (4.17) and satisfies the monotonicity constraints (4.18) with strict inequality (recall,  $\alpha_1 < \dots < \alpha_k$ ).<sup>7</sup> Then, expressing each allocation  $e_t$  in terms of the single parameter  $\delta$  and substituting into the  $t$ -th participation constraint (4.17) yields

$$\delta \geq \frac{U^0(\frac{t}{k})}{\sum_{s=1}^t \frac{1}{k-s}},$$

where  $\alpha_t$  has been replaced with  $t/k$  to simplify the right-hand side expression. To minimize  $\delta$  subject to the participation constraints, it must therefore be that

$$\delta = \frac{U^0(\frac{t^*}{k})}{\sum_{s=1}^{t^*} \frac{1}{k-s}}, \quad \text{where } t^* \in \arg \max_{t \in T} \frac{U^0(\frac{t}{k})}{\sum_{s=1}^t \frac{1}{k-s}}. \quad (4.20)$$

The value of (4.15–4.19), and hence (4.10–4.14), is thus given by (4.20).

To find the random linear contract that attains the guarantee, we can use the complementary slackness conditions, together with (4.11) and (4.12). Because all monotonicity constraints (4.18) hold with strict inequality, complementary slackness implies that  $\theta_t = 0$  for all  $t$  in the corresponding solution to (4.10–4.14). Moreover, (4.17) holds with equality only for  $t = t^*$ . Thus, by complementary slackness,  $\lambda_t = 0$  for all  $t \neq t^*$ . Because (4.11) must bind for each  $t$ , it follows that  $p_t = 0$  for all  $t > t^*$ . Moreover, for  $t \leq t^*$ , (4.11) yields

$$p_t = \frac{\lambda_{t^*}}{k-t},$$

again using  $\alpha_t = t/k$ . Substituting each  $p_t$  for the right-hand side expression and using the probability constraint (4.12) pins down the exact value of  $\lambda_{t^*}$ . It follows from simple arithmetic that the optimal grid-based contract sets

$$p_t = \begin{cases} \frac{1}{k-t} \left( \sum_{s=1}^{t^*} \frac{1}{k-s} \right)^{-1} & \text{if } 1 \leq t \leq t^*, \\ 0 & \text{if } t^* < t \leq k. \end{cases} \quad (4.21)$$

It can be shown that as the number of grid points  $k$  grows to infinity, (4.20) converges to (5.1), the value proven to be the optimal guarantee. When there exists a positive solution to the maximization problem in (5.1), then the limit of (4.21) is the random linear contract corresponding to the cumulative distribution function satisfying (5.2), the contract proven to be optimal. Otherwise, the sequence of solutions converges to the point mass on the zero

---

<sup>7</sup>Because  $\alpha_k = 1$ , we can take  $e_k = e_{k-1} + 1$  without loss of optimality.

contract. We show below that this is precisely the corner case in which randomization does not have value (see Corollary 3).

## 5 Saddle point

To identify the optimal guarantee and an optimal contract, we construct a saddle point in the zero-sum game played between the principal and Nature using our guess from Section 4.

**Proposition 3** (Optimal Guarantee and Contract).

1. *The optimal guarantee is<sup>8</sup>*

$$\bar{v} := \max_{\alpha \in [0,1]} \frac{U^0(\alpha)}{-\ln(1-\alpha)} > 0. \quad (5.1)$$

2. *If there exists  $\alpha^* > 0$  attaining the maximum in (5.1), then the optimal guarantee is achieved by the random linear contract  $p^* \in \bar{\Delta}([0,1])$  corresponding to the cumulative distribution function  $G_{p^*} : \mathbb{R} \rightarrow [0,1]$  defined by*

$$G_{p^*}(\alpha) := \frac{\ln(1-\alpha)}{\ln(1-\alpha^*)} \quad \text{for } \alpha \in [0, \alpha^*]. \quad (5.2)$$

*Otherwise, the optimal guarantee is attained in the limit of a sequence of deterministic linear contracts converging to the zero contract.*

Part 2 of Proposition 3 gives the distribution of the agent's share  $\alpha$  under the optimal random linear contract  $p^*$ . Simple algebra shows that the corresponding cumulative distribution function for the principal's share  $1 - \alpha$  is given by

$$H_{p^*}(1 - \alpha) := 1 - \frac{\ln(1 - \alpha)}{\ln(1 - \alpha^*)} \quad \text{for } 1 - \alpha \in [1 - \alpha^*, 1].$$

That is, the logarithm of the principal's share, or  $\ln(1 - \alpha)$ , is distributed uniformly on  $[\ln(1 - \alpha^*), 0]$ . In other words, the principal's share has a log-uniform distribution, also known as the reciprocal distribution.

---

<sup>8</sup>By convention, we extend the quotient to all of  $[0,1]$  by left and right limits. That is, let

$$\frac{U^0(0)}{-\ln(1-0)} := \lim_{\alpha \rightarrow 0^+} \frac{U^0(\alpha)}{-\ln(1-\alpha)} \in [-\infty, \infty) \quad \text{and} \quad \frac{U^0(1)}{-\ln(1-1)} := \lim_{\alpha \rightarrow 1^-} \frac{U^0(\alpha)}{-\ln(1-\alpha)} = 0,$$

where the first limit is bounded from above (because  $U^0(0) \leq 0$ ), ensuring the existence of a maximum in (5.1). The maximum is positive because  $U^0(\alpha) > 0$  for  $\alpha$  sufficiently close to one by the existence of a known surplus-generating action. Hence, any maximizer must be strictly smaller than one.

The rest of this section is dedicated to the proof of Proposition 3.

## 5.1 Attaining $\bar{v}$

We show that the conjectured optimal contract attains  $\bar{v}$ , defined in (5.1) in Proposition 3. Suppose first that there exists  $\alpha^* \in (0, 1)$  attaining the maximum in (5.1). Let  $p^*$  be the random linear contract corresponding to (5.2). We show that  $p^*$  obtains a guarantee of at least  $\bar{v}$ .

By applying the Revelation Principle, we again formulate Nature's problem of choosing a technology  $A$  to minimize the profit  $V(p, A)$  as a mechanism design problem where the realized contract  $\alpha$  is the agent's type, the expected output  $e$  is the allocation, the cost  $c$  plays the role of a transfer, and the agent's outside option is to play a known action. For any random linear contract  $p$ , let

$$L(p) := \min_{e(\cdot), c(\cdot)} \int_0^1 (1 - \alpha) e(\alpha) dG_p(\alpha) \quad \text{s.t.} \quad (5.3)$$

$$\alpha e(\alpha) - c(\alpha) \geq \alpha e(\alpha') - c(\alpha') \quad \forall \alpha, \alpha' \in [0, 1] : \alpha \neq \alpha', \quad (5.4)$$

$$\alpha e(\alpha) - c(\alpha) \geq U^0(\alpha) \quad \forall \alpha \in [0, 1], \quad (5.5)$$

$$c(\alpha) \geq 0, \quad 0 \leq e(\alpha) \leq y_n \quad \forall \alpha \in [0, 1]. \quad (5.6)$$

Note that the minimization problem (5.3–5.6) is well-defined and the existence of a minimum follows by general existence results for principal-agent problems with adverse selection (e.g., Nöldeke and Samuelson (2018), Proposition 9). It is also worth noting that constraints (5.4–5.6) are taken to hold for all  $\alpha, \alpha' \in [0, 1]$ , not just for  $\alpha, \alpha' \in \text{supp}(p)$ . This is convenient as it allows us to work with functions  $e(\cdot)$  and  $c(\cdot)$  defined on the full interval  $[0, 1]$  even when the support of  $p$  is only a subset thereof.

The following Lemma uses standard arguments to simplify constraints (5.4–5.6). It also verifies that  $L(p)$  is, in fact, the principal's payoff guarantee, taking into account the requirement that Nature's technology must be compact and the assumption of adversarial tie-breaking.

**Lemma 4.** *For any random linear contract  $p$ ,*

$$V(p) = L(p) = \min_{e(\cdot)} \int_0^1 (1 - \alpha) e(\alpha) dG_p(\alpha) \quad s.t. \quad (5.7)$$

$$e(\cdot) \text{ is nondecreasing}, \quad (5.8)$$

$$\int_0^\alpha e(t) dt \geq U^0(\alpha) \quad \forall \alpha \in [0, 1], \quad (5.9)$$

$$e(0) \geq 0, \quad e(1) \leq y_n. \quad (5.10)$$

Now, suppose  $e^*(\cdot)$  is a feasible solution to (5.8–5.10) and attains the minimum in (5.3). By Lemma 4 and the definition of  $G_{p^*}$ , the guarantee of  $p^*$  is given by

$$\begin{aligned} V(p^*) &= \int_0^1 (1 - \alpha) e^*(\alpha) dG_{p^*}(\alpha) = \int_0^{\alpha^*} (1 - \alpha) e^*(\alpha) d\left(\frac{\ln(1 - \alpha)}{\ln(1 - \alpha^*)}\right) \\ &= \int_0^{\alpha^*} \frac{e^*(\alpha)}{-\ln(1 - \alpha^*)} d\alpha = \frac{\int_0^{\alpha^*} e^*(\alpha) d\alpha}{-\ln(1 - \alpha^*)} \geq \frac{U^0(\alpha^*)}{-\ln(1 - \alpha^*)} = \bar{v}, \end{aligned}$$

where the inequality is from (5.9). This establishes the desired result.

We now consider the corner case in which the unique maximizer of (5.1) is  $\alpha^* = 0$ . We show that  $\bar{v}$  is attained in the limit of a sequence of deterministic contracts converging to the zero contract. From the formula for the deterministic guarantee given in (A.2), this is equivalent to showing that

$$\lim_{\alpha \rightarrow 0^+} \frac{1 - \alpha}{\alpha} U^0(\alpha) \geq \bar{v}.$$

Because the left-hand side expression is larger than zero (by (A.1)), the inequality holds if  $\bar{v} < 0$ . If  $\bar{v} \geq 0$ , then it must be that  $U^0(0) = 0$ . Hence, the known technology  $A^0$  contains a zero cost action. Let  $e^0$  denote the maximum expected output among actions with zero cost. Suppose  $(\alpha_k)_k$  is a positive sequence that converges to zero. Take any corresponding sequence of (principal least-preferred) best responses in  $A^0$ ,  $(e(\alpha_k), c(\alpha_k))_k$ . Then,

$$\alpha_k e^0 \leq \alpha_k e(\alpha_k) - c(\alpha_k) \iff c(\alpha_k) \leq \alpha_k (e(\alpha_k) - e^0) \leq \alpha_k y_n,$$

where the last inequality follows because output is bounded below by 0 and above by  $y_n$ . Hence, as  $k \rightarrow \infty$ , it must be that  $c(\alpha_k) \rightarrow 0$ . Because the tail of  $(e(\alpha_k), c(\alpha_k))_k$  belongs to a compact set  $[e^0, y_n] \times [0, \bar{c}]$  for some  $\bar{c} > 0$ , we may extract a convergent subsequence  $(\alpha_{k_\ell})_{k_\ell}$  with  $e(\alpha_{k_\ell}) \rightarrow \bar{e}$  and  $c(\alpha_{k_\ell}) \rightarrow 0$ . Thus,

$$\bar{v} = \lim_{\alpha \rightarrow 0^+} \frac{U^0(\alpha)}{-\ln(1 - \alpha)} = \lim_{k_\ell \rightarrow \infty} \frac{U^0(\alpha_{k_\ell})}{-\ln(1 - \alpha_{k_\ell})} = \lim_{k_\ell \rightarrow \infty} \frac{\alpha_{k_\ell} e(\alpha_{k_\ell})}{-\ln(1 - \alpha_{k_\ell})} = \bar{e}.$$

On the other hand,

$$\lim_{\alpha \rightarrow 0^+} \frac{1-\alpha}{\alpha} U^0(\alpha) = \lim_{k_\ell \rightarrow \infty} \frac{1-\alpha_{k_\ell}}{\alpha_{k_\ell}} U^0(\alpha_{k_\ell}) = \lim_{k_\ell \rightarrow \infty} \frac{1-\alpha_{k_\ell}}{\alpha_{k_\ell}} (\alpha_{k_\ell} e(\alpha_{k_\ell}) - c(\alpha_{k_\ell})) = \bar{e} \geq \bar{v}.$$

## 5.2 $\bar{v}$ is an upper bound

We conclude the proof of Proposition 3 by showing that no contract can obtain a guarantee higher than  $\bar{v}$ , again defined in (5.1) in Proposition 3. It suffices to exhibit a technology  $A \supseteq A^0$  such that no *deterministic* contract can attain a payoff higher than  $\bar{v}$  against  $A$  (no matter how ties are broken among optimal actions for the agent). For the purposes of constructing such a technology, let  $\alpha^* \in [0, 1)$  attain the maximum in (5.1) (by footnote 8,  $\alpha^* < 1$ ). Moreover, let  $e^*$  be given by

$$e^*(\alpha) := \min \left\{ y_n, \frac{U^0(\alpha^*)}{-(1-\alpha) \ln(1-\alpha^*)} \right\}, \quad \forall \alpha \in [0, 1], \quad (5.11)$$

and let  $c^*$  be given by

$$c^*(\alpha) := \alpha e^*(\alpha) - \int_0^\alpha e^*(t) dt, \quad \forall \alpha \in [0, 1].$$

Finally, let  $\bar{\alpha}$  satisfy  $e^*(\bar{\alpha}) = \max_{(\pi_0, c_0) \in A^0} \pi_0 y$ .<sup>9</sup> (To see where (5.11) comes from, observe that it corresponds to making the dual constraint (4.16) bind, save for the upper bound on output.)

We now identify a worst-case action set. For every  $\alpha \in [0, \bar{\alpha}]$ , let

$$\begin{aligned} A(0) &:= \{(\pi, c) | \pi y \leq e^*(0), \max\{y_n, \bar{c}\} \geq c \geq 0\}, \\ A(\alpha) &:= \{(\pi, c) | \pi y = e^*(\alpha), \max\{y_n, \bar{c}\} \geq c \geq c^*(\alpha)\} \quad \forall \alpha \in (0, \bar{\alpha}], \end{aligned}$$

where  $\bar{c}$  denotes the maximum cost of any action in  $A^0$ . By construction,  $A := \cup_{\alpha \in [0, \bar{\alpha}]} A(\alpha)$  is compact, i.e., a technology. We show, in addition, that it contains  $A^0$ .

**Lemma 5.**  $A \supseteq A^0$ .

*Proof.* For an arbitrary  $(\pi_0, c_0) \in A^0$ , if  $\pi_0 y \leq e^*(0)$ , we have that  $(\pi_0, c_0) \in A(0)$ . Otherwise,

---

<sup>9</sup>Note that such an  $\bar{\alpha}$  exists because  $\max_{(\pi_0, c_0) \in A^0} \pi_0 y \leq y_n = e^*(1)$ . Moreover,

$$e^*(0) \leq \frac{U^0(\alpha^*)}{-\ln(1-\alpha^*)} \leq \lim_{\alpha \rightarrow \alpha^*} \frac{\alpha (\max_{(\pi_0, c_0) \in A^0} \pi_0 y)}{-\ln(1-\alpha)} \leq \max_{(\pi_0, c_0) \in A^0} \pi_0 y.$$

The final inequality follows because  $\alpha^* \leq -\ln(1-\alpha^*)$  for  $\alpha^* \in (0, 1)$  and  $\lim_{\alpha \rightarrow 0} (-\alpha/\ln(1-\alpha)) = 1$ .

let  $\pi_0 y = e^*(\alpha_0)$  for some  $\alpha_0 \in (0, \bar{\alpha}]$ . We then have

$$\begin{aligned} \alpha_0 e^*(\alpha_0) - c^*(\alpha_0) &= \int_0^{\alpha_0} e^*(t) dt = \int_0^{\alpha_0} \frac{U^0(\alpha^*)}{-(1-t) \ln(1-\alpha^*)} dt = \frac{-\ln(1-\alpha_0) U^0(\alpha^*)}{-\ln(1-\alpha^*)} \\ &\geq U^0(\alpha_0) = \max_{(\pi, c) \in A^0} (\alpha_0 \pi y - c) \geq \alpha_0 \pi_0 y - c_0 = \alpha_0 e^*(\alpha_0) - c_0, \end{aligned}$$

where the first inequality follows from the definition of  $\alpha^*$ . Consequently,  $\bar{c} \geq c_0 \geq c^*(\alpha_0)$  and  $(\pi_0, c_0) \in A(\alpha_0)$ .  $\square$

For any (potentially non-linear) deterministic contract  $w$ , we show that

$$V(w, A) \leq \bar{v}.$$

Let  $a_w = (\pi_w, c_w) \in A$  be a best response of the agent. Then,  $a_w$  maximizes the agent's utility, i.e.,

$$\pi_w w - c_w = \max_{(\pi, c) \in A} (\pi w - c).$$

Suppose  $a_w \in A(\alpha_w)$ . Then  $\pi_w y \leq e^*(\alpha_w)$  and  $c_w = c^*(\alpha_w)$ . Consider the following subset of  $A$ :

$$\left\{ (t\pi_w + (1-t)0, c^*((e^*)^{-1}(te^*(\alpha_w)))) \right\}_{t \in [0,1]},$$

where  $0$  is the  $n$ -dimensional zero vector and, by convention,  $(e^*)^{-1}(y_n) := 1$ . For each  $t \in [0, 1]$ , the agent's corresponding utility is

$$u_w(t) := t\pi_w w - c^*((e^*)^{-1}(te^*(\alpha_w))).$$

Observe that  $t = 1$  corresponds to the agent's best action. Hence, taking  $u'_w(1)$  to be the left-derivative of  $u_w$  at 1,<sup>10</sup>

$$\begin{aligned} 0 \leq u'_w(1) &= \pi_w w - (c^*)'((e^*)^{-1}(te^*(\alpha_w))) \left( (e^*)^{-1} \right)'(te^*(\alpha_w)) e^*(\alpha_w) \Big|_{t=1} \\ &= \pi_w w - (c^*)'(\alpha_w) \frac{1}{(e^*)'(\alpha_w)} e^*(\alpha_w) \\ &= \pi_w w - \alpha_w e^*(\alpha_w). \end{aligned}$$

Here, the first equality follows from the chain rule; the second equality follows from the inverse function theorem; and the last equality follows from the definition of  $c^*$ . Notice that

$$0 \leq \pi_w w - \alpha_w e^*(\alpha_w) \iff \alpha_w e^*(\alpha_w) \leq \pi_w w.$$

---

<sup>10</sup>The left-derivative is well-defined because  $(e^*)^{-1}(\cdot)$  is a differentiable bijection on  $(0, 1)$ .

Consequently, the expected payoff of the principal is at most

$$\pi_w(y - w) \leq (1 - \alpha_w)e^*(\alpha_w) = \frac{U^0(\alpha^*)}{-\ln(1 - \alpha^*)} = \bar{v}.$$

This concludes the proof.

## 6 Gain from randomization

[Kambhampati \(2023\)](#) showed that a sufficient condition for the principal to benefit from randomization is that there is an optimal deterministic contract different from the zero contract. We provide here a weaker necessary and sufficient condition.

Define the *gain from randomization* as the ratio between the optimal guarantee and the optimal deterministic guarantee, or  $\sup_{p \in \Delta(\mathbb{R}_+^n)} V(p) / \sup_{w \in \mathbb{R}_+^n} V(w)$ .<sup>11</sup> We say that the gain is *positive* if the ratio is strictly greater than 1.

**Corollary 3.** *The gain from randomization is positive if and only if*<sup>12</sup>

$$\frac{U^0(0)}{-\ln(1 - 0)} < \max_{\alpha \in [0,1]} \frac{U^0(\alpha)}{-\ln(1 - \alpha)}.$$

When the inequality in Corollary 3 is violated, the optimal guarantee is obtained in the limit of a sequence of deterministic linear contracts converging to the zero contract. Hence, there is no value to randomization. When the inequality holds, any optimal contract has a non-degenerate support and randomization benefits the principal. The following example shows that this can be the case even when the optimal deterministic contract is (essentially) the zero contract. In this case, randomization results in a pure efficiency gain that benefits both the principal and the agent. This provides a stark contrast to the rent extraction argument of [Kambhampati \(2023\)](#), who shows that randomization has value when the optimal

---

<sup>11</sup>To see why this is the relevant object of interest, note that the principal's utility is measured on some von Neumann–Morgenstern scale. The utility difference between contracting with the agent using the optimal randomized contract and not contracting with the agent at all is  $(\sup_{p \in \Delta(\mathbb{R}_+^n)} V(p) + v_0) - v_0 = \sup_{p \in \Delta(\mathbb{R}_+^n)} V(p)$ , where  $v_0$  is the principal's utility from sources other than the contracting problem at hand. Similarly,  $\sup_{w \in \mathbb{R}_+^n} V(w)$  is the utility difference between contracting with the agent using the optimal deterministic contract and not contracting with the agent at all. Thus,  $\sup_{p \in \Delta(\mathbb{R}_+^n)} V(p) / \sup_{w \in \mathbb{R}_+^n} V(w)$  is a ratio of utility differences, and hence invariant to the choice of scale (i.e., to positive affine transformations of utility), making its magnitude unit-free and meaningful. In contrast, the difference  $\sup_{p \in \Delta(\mathbb{R}_+^n)} V(p) - \sup_{w \in \mathbb{R}_+^n} V(w)$  depends on the choice of units and can be scaled arbitrarily when different from zero.

<sup>12</sup>Recall our convention in footnote 8 to treat the maximand of the guarantee as an extended real-valued function defined on  $[0, 1]$ .

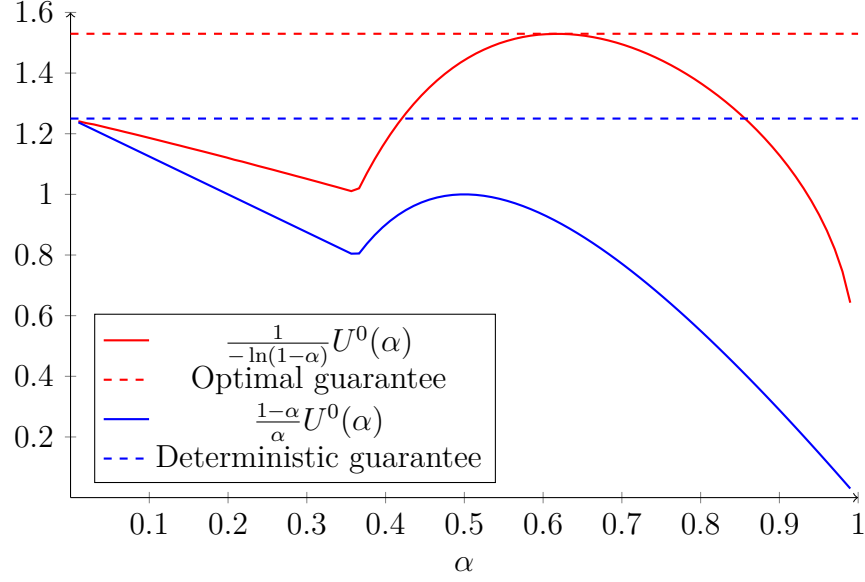


Figure 6.1. Optimal guarantee versus deterministic guarantee.

deterministic linear contract is non-zero because it can reduce payments to the agent while holding worst-case productivity fixed.

**Example 1.** Suppose the known technology  $A^0$  consists of two actions,  $(\pi, 1)$  and  $(\pi', 0)$ , with expected outputs  $\pi y = 4$  and  $\pi' y = 5/4$ . Then,  $U^0(\alpha) = \max\{4\alpha - 1, 5\alpha/4\}$  and the optimal deterministic guarantee from [Carroll \(2015\)](#) is  $\max_{\alpha \in [0,1]} \frac{1-\alpha}{\alpha} U^0(\alpha) = 5/4$ . Under principal-optimal tie-breaking, the guarantee would be obtained by the zero contract, which “targets” the zero-cost action  $(\pi', 0)$ . Because of adversarial tie-breaking, however, it is here obtained in the limit of a sequence of linear contracts converging to the zero contract (see [Appendix A](#), which reproduces the deterministic analysis under adversarial tie-breaking). [Figure 6.1](#) provides an illustration.

On the other hand, the optimal guarantee is 1.530, which is strictly larger than the optimal deterministic guarantee of 5/4. The optimal guarantee is obtained by setting the upper bound of the support of the optimal random contract [\(5.2\)](#) to  $\alpha^* \approx 0.618$ . The basic intuition for the gain from randomization is that it allows the principal to “target” the positive-cost known action  $(\pi, 1)$  to increase efficiency, while extracting enough rent from the agent to make this worthwhile.

Notice that the agent also benefits from the randomization. Rather than being offered (essentially) the zero contract, he is (almost surely) given a contract with positive slope and thus, no matter the true technology, can obtain a positive payoff from the costless known action  $(\pi', 0)$ .

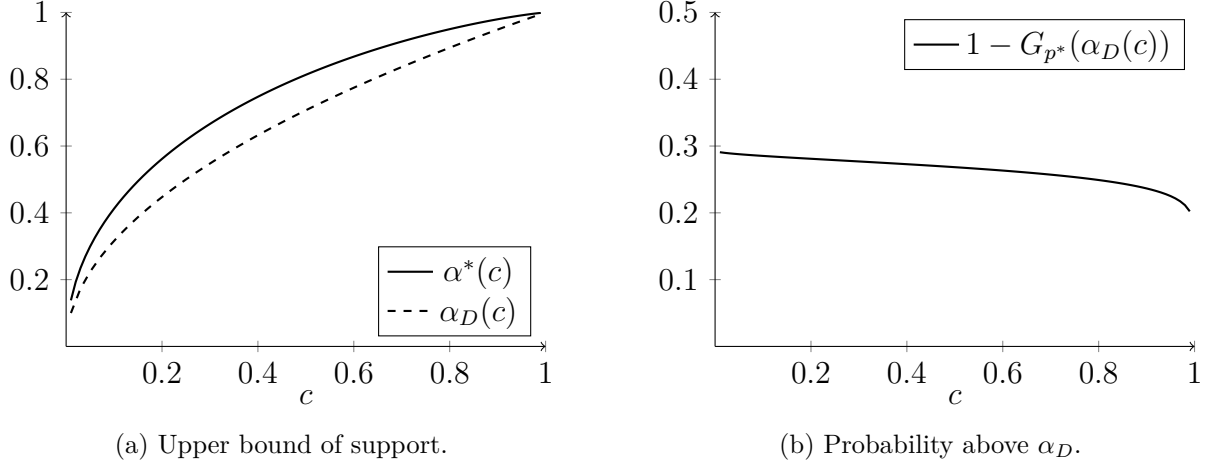


Figure 7.1. Comparison of the optimal random and deterministic linear contracts.

## 7 Comparative statics under maximal uncertainty

We now derive comparative statics of the optimal contract and guarantee under maximal uncertainty about the agent's action set, i.e., when there is only one known action. For example, the principal may know of a single way for the agent to approach a problem, but be uncertain of all other possible ways to approach it. Upon normalizing the expected output of the known action, the predictions in this case are summarized by a single parameter: the effort cost of the known action.

**Proposition 4.** *Suppose the known technology  $A^0$  consists of a single action  $(\pi, c)$ , whose expected output  $\pi y$  is normalized to 1 and effort cost is  $c \in (0, 1)$ . Then the following properties hold:*

1. *The upper bound of the support of the optimal random linear contract,  $\alpha^*(c)$ , is larger than the slope of the optimal deterministic linear contract,  $\alpha_D(c)$ . Moreover,  $\alpha^*(c)$  is increasing and concave in  $c$  with  $\lim_{c \rightarrow 0} \alpha^*(c) = 0$  and  $\lim_{c \rightarrow 1} \alpha^*(c) = 1$ .*
2. *The optimal guarantee,  $\bar{v}(c)$ , is decreasing and convex in  $c$  with  $\lim_{c \rightarrow 0} \bar{v}(c) = 1$  and  $\lim_{c \rightarrow 1} \bar{v}(c) = 0$ .*
3. *The gain from randomization,  $\gamma(c)$ , satisfies  $\lim_{c \rightarrow 0} \gamma(c) = 1$  and  $\lim_{c \rightarrow 1} \gamma(c) = +\infty$ .*

We further illustrate Proposition 4 via a sequence of figures. Figure 7.1a illustrates the relationship between the upper bound of the support of the optimal random linear contract and the slope of the optimal deterministic contract as a function of the effort cost. Figure 7.1b depicts the probability that the agent receives a strictly higher share of output than under the optimal deterministic contract. Inspecting the figure, we see that

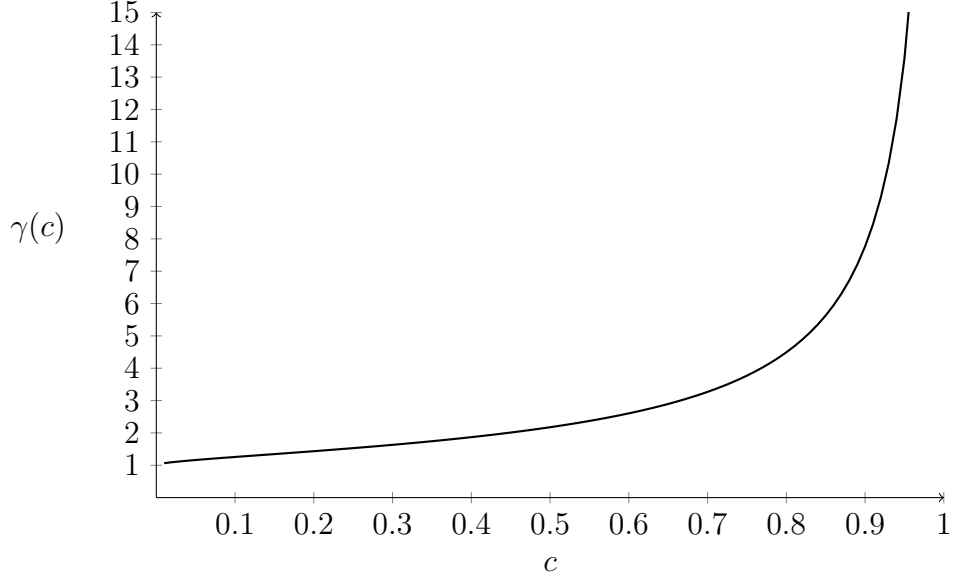


Figure 7.2. Gain from randomization.

this probability is substantive for a wide range of effort costs. For instance, at  $c = 0.5$ , it is about 27%. Whenever this happens, the agent takes a (weakly) more productive action than under the optimal deterministic contract no matter the true technology. Thus, the gain from randomization comprises both efficiency gains (cf. Example 1) and rent extraction (cf. Kambhampati, 2023). Finally, Figure 7.2 shows that, as the cost of effort grows and the moral hazard problem becomes more severe, the gain from randomization grows arbitrarily large. In the other direction, as the cost of effort approaches zero and the optimal random linear contract weakly converges to the point mass on the zero contract, the gain from randomization approaches its minimal value of one.

## 8 Discussion

We have shown that linear contracts remain robustly optimal when randomization is allowed. We have also identified an optimal contract where the principal's share has a log-uniform distribution. Though we have focused attention on the Carroll (2015) model, we suspect both that randomization improves the principal's guarantee and that our techniques will be useful in characterizing optimal contracts in related models.<sup>13</sup>

We conclude by discussing corollaries and extensions of our results.

---

<sup>13</sup>E.g., the working paper by Peng and Tang (2024) considered the team model of Dai and Toikka (2022).

## 8.1 Commitment

The proof of Proposition 3 constructs a saddle-point  $(p, A)$  for the principal’s max-min problem, i.e.,  $p$  is a best response for the principal against  $A$  and  $A$  is a best response for Nature against  $p$ . Hence, given the worst-case technology  $A$ , the principal is indifferent among all contracts in the support of  $p$ . It follows that commitment to randomization is not needed. That is, the principal has no incentive to deviate from any of her realized contracts.

## 8.2 Tie-breaking

We have assumed that, when indifferent among actions, the agent chooses the one least preferred by the principal. However, the value (5.1) remains the optimal guarantee if instead we assume principal-optimal tie-breaking. This follows because the guarantee of the optimal contract (5.2) can only increase and the proof given in Section 5.2 shows that the value (5.1) remains an upper bound. The only modification required in the statement of Proposition 3 is in the corner case in which randomization has no value. With principal-optimal tie-breaking, the optimal guarantee is attained by the zero contract with no need for approximation.

## 8.3 Screening

Because the agent knows the true technology, it is natural to ask whether the principal’s guarantee could be increased by using more general contracts that elicit this information, or equivalently, by offering the agent a menu of contracts. It follows immediately from Proposition 3 that this is not the case. The performance of any menu of contracts is no better than its performance against the worst-case technology  $A$  constructed in the proof of Proposition 3. Specifically, against  $A$ , no matter which contract the agent chooses from the menu, the resulting payoff to the principal is no higher than  $V(p)$ . It follows that randomizing over menus cannot improve the principal’s guarantee.

## 8.4 Extensions

We have deliberately kept the model streamlined to allow for a simple and transparent analysis of the merits of linear contracts, but we note here three extensions that are immediate.

First, the improvement argument can accommodate the introduction of a participation constraint with minor modifications. The simplest way to impose a participation constraint is to assume that the known technology  $A^0$  contains an action that results in zero output at no cost to the agent—this is subsumed by our general analysis. If instead we assume that any contract has to deliver the agent an expected payoff (net of cost) of at least  $\bar{u} > 0$ , then this

can be handled by replacing  $U^0(w)$  with  $\max\{U^0(w), \bar{u}\}$  on the right-hand side of (3.12). Because the affine contract  $\hat{w}^\tau$  constructed in the proof of Lemma 1 pointwise increases the corresponding original contract  $w^\tau$ , it satisfies the so-modified participation constraint (3.12). Moreover, it is straightforward to verify that  $\hat{w}^\tau$  is affine as long as all its in-neighbors are affine. We thus recover a modified version of Lemma 1 where the assumption is weakened to requiring all in-neighbors of  $\tau$  to be affine and the conclusion is changed to asserting the existence of an affine contract and a corresponding random contract that improves on the original one. The rest of the argument proceeds as before to iteratively yield, for any finitely supported random contract, a random affine contract that (weakly) outperforms it. Thus, random affine contracts are optimal as a class.

Second, it is straightforward to allow for multi-dimensional output  $y$  from any finite set  $Y$ , with the payoff to the principal being  $v(y)$  for some function  $v : Y \rightarrow \mathbb{R}_+$ . Modifying the primal problem (3.10–3.14) in the obvious way, the same argument gives an improvement from moving to a random contract that is linear in  $v(y)$ . Specifically, replace each  $y_i$  with  $v(y_i)$  in the primal objective (3.10) and on the right-hand side of the dual constraint (3.16).

Finally, a compact output set  $Y \subset \mathbb{R}_+$  would substantially complicate the duality-based arguments as Nature’s primal and dual problems would then no longer be finite. However, the proof of Proposition 3 goes through verbatim. Thus, our characterization of the optimal guarantee and of the optimal random linear contract carries over as stated to a compact  $Y$ .

## Appendix A: Deterministic analysis

Consider Nature’s problem (3.1–3.4) given a linear contract with slope  $\alpha$ :

$$V(\alpha) = \min_{\pi, c} (1 - \alpha)\pi y \quad \text{s.t.} \quad \alpha\pi y - c \geq U^0(\alpha), \quad \sum_{i=1}^n \pi_i = 1, \quad \text{and} \quad c, \pi_i \geq 0 \quad \forall i \in I.$$

It is clearly optimal to set  $c = 0$  and choose  $\pi$  such that  $\alpha\pi y = U^0(\alpha)^+$ , where our convention is to write  $b^+ = \max\{b, 0\}$  for any  $b \in \mathbb{R}$ . If  $\alpha = 0$ , then it is optimal to put  $\pi y = U^0(0)^+ = 0$ , and thus the guarantee from the zero contract is zero. (This is of course immediate given adversarial tie-breaking.) On the other hand, the guarantee for any positive slope  $\alpha > 0$  is  $V(\alpha) = (1 - \alpha)U^0(\alpha)^+/\alpha$ . Recalling the definition of  $U^0$  from (2.2), we thus have

$$V(\alpha) = \begin{cases} 0 & \text{if } \alpha = 0, \\ (1 - \alpha) \max_{(\pi, c) \in A^0} \left( \pi y - \frac{c}{\alpha} \right)^+ & \text{if } \alpha \in (0, 1]. \end{cases} \quad (\text{A.1})$$

By Corollary 1, the optimal deterministic guarantee is given by the supremum of the

function  $V : [0, 1] \rightarrow \mathbb{R}_+$  defined by (A.1). The non-triviality assumption implies that  $V(\alpha) > 0$  for  $\alpha$  sufficiently close to 1, and thus the optimal guarantee is positive. It can also readily be verified to coincide with Carroll’s (2015) guarantee for principal-optimal tie-breaking.<sup>14</sup> In what follows, we will use the fact that the optimal deterministic guarantee can also be written as

$$\sup_{w \in \mathbb{R}_+^n} V(w) = \max_{\alpha \in [0, 1]} \frac{1 - \alpha}{\alpha} U^0(\alpha), \quad (\text{A.2})$$

provided that  $\frac{(1-\alpha)}{\alpha} U^0(\alpha)$  is taken to mean the limit as  $\alpha \rightarrow 0$ .

An optimal deterministic contract, which is linear, can be found by simply maximizing the guarantee (A.1). However, there is a small wrinkle as  $V$ , which is continuous on  $(0, 1]$ , may have a downward jump in the limit as  $\alpha \rightarrow 0$ . This leads to the following characterization.<sup>15</sup>

**Proposition A.1.** *A linear contract  $\alpha$  is an optimal deterministic contract if and only if it maximizes the function  $V : [0, 1] \rightarrow \mathbb{R}_+$  defined by (A.1). (Because of the existence of a known surplus-generating action, any maximizer must be interior.) If no such maximizer exists, then there does not exist an optimal deterministic contract. In this case, the optimal deterministic guarantee is attained in the limit of any sequence of linear contracts  $\alpha_n > 0$  such that  $\alpha_n \rightarrow 0$ .*

*Proof.* The first claim characterizing optimal linear contracts follows by Corollary 1 and the construction of  $V$ . That there is no optimal deterministic contract when there is no linear optimal contract follows by Proposition 1. Finally, the only possible point of discontinuity of  $V$  is 0, and hence if a maximizer does not exist, then  $\sup_{\alpha \in [0, 1]} V(\alpha) = \lim_{\alpha \rightarrow 0^+} V(\alpha)$ .  $\square$

While adversarial tie-breaking generates an existence problem relative to the case of principal-optimal tie-breaking where an optimal (linear) contract always exists, the issue is arguably minor: It only arises in the corner case where under principal-optimal tie-breaking, the zero contract is the uniquely optimal linear contract. This case seems uninteresting from the perspective of optimal incentive provision as none is required. It is incidentally also the only case where randomization does not improve the principal’s payoff—see Corollary 3. A simple sufficient condition to rule it out and to ensure the existence of an optimal contract is for any known action generating a positive surplus to have a non-zero cost.

<sup>14</sup>Principal-optimal tie-breaking leads in general to a weakly higher guarantee from the zero contract than adversarial tie-breaking, but the guarantees are the same for any positive slope. As the guarantee under principal-optimal tie-breaking is continuous in  $\alpha$ , this implies that the optimal guarantees are the same.

<sup>15</sup>It is easy to show that all optimal deterministic contracts are linear under Carroll’s (2015) full support condition by verifying that then certain inequalities in the proof of Proposition 1 are strict. As this argument is similar to Carroll’s, we omit it in the interest of space.

## Appendix B: Proofs

### B.1 Proof of Lemma 2

We first prove a stronger result for an  $\varepsilon$ -perturbation of the problem and then establish Lemma 2 via taking the limit  $\varepsilon \rightarrow 0$ .

Given  $\varepsilon \geq 0$ , we introduce a relaxation of the incentive compatibility constraint (3.11):

$$\pi^t w^t - c_t \geq \pi^s w^t - c_s - \varepsilon \quad \forall t, s \in T : t \neq s. \quad (\text{B.1})$$

The  $\varepsilon$ -perturbed primal problem consists of solving (3.10) subject to (B.1), (3.12), (3.13), and (3.14). It is feasible for all  $\varepsilon \geq 0$ , because it is a relaxation of the feasible primal problem (3.10–3.14). Let  $V(p; \varepsilon)$  denote the value of the  $\varepsilon$ -perturbed primal problem. We clearly have  $V(p; \varepsilon) \geq -\max_{i,t} w_i^t$ , and hence an optimal solution exists.

Because the perturbation  $\varepsilon$  only enters the right-hand sides of some primal constraints, in the dual it only enters the objective. The  $\varepsilon$ -perturbed dual problem consists of solving

$$V(p; \varepsilon) = \max_{\kappa, \lambda, \mu} \sum_{t \in T} (\lambda_t U(w^t) + \mu_t) - \varepsilon \sum_{t \in T} \sum_{s \neq t} \kappa_{st} \quad (\text{B.2})$$

subject to (3.16–3.18).

**Lemma B.1.** *Let  $\varepsilon > 0$ . If  $(\kappa^*, \lambda^*, \mu^*)$  is an optimal solution to the  $\varepsilon$ -perturbed dual problem, then  $G(\kappa^*)$  is acyclic.*

*Proof.* Let  $\varepsilon > 0$  and let  $(\kappa^*, \lambda^*, \mu^*)$  be an optimal solution to the  $\varepsilon$ -perturbed dual. Suppose toward contradiction that  $G(\kappa^*)$  contains a cycle. By relabeling if necessary, we can assume without loss of generality that the cycle involves vertices  $C := \{1, \dots, m\}$  for  $2 \leq m \leq k$ , with  $\kappa_{t,t+1}^* > 0$  for all  $t \in C$ , where  $m+1 = 1$  by convention. Similarly, we can assume without loss of generality that  $p_1 = \min_{t \in C} p_t$ .

By complementary slackness, the incentive compatibility constraints in (B.1) corresponding to  $\kappa_{t,t+1}^*$  for  $t \in C$  hold with equality. Letting  $\{(\pi^t, c_t)\}_{t \in T}$  denote an optimal solution to the  $\varepsilon$ -perturbed primal, we thus have

$$\pi^t w^t - c_t - \pi^{t+1} w^t + c_{t+1} = -\varepsilon \quad \forall t \in C. \quad (\text{B.3})$$

Summing the equalities over  $t \in C$  and recalling that  $m+1 = 1$  gives

$$\sum_{t \in C} (\pi^t - \pi^{t+1}) w^t = -m\varepsilon. \quad (\text{B.4})$$

We will show that it is then possible to strictly lower the value of the  $\varepsilon$ -perturbed primal, contradicting optimality.

Construct a new solution for the  $\varepsilon$ -perturbed primal as follows. Let  $(\hat{\pi}^t, \hat{c}_t) := (\pi^t, c_t)$  for all  $t \notin C$ . For each  $t \in C$ , let

$$(\hat{\pi}^t, \hat{c}_t) := (1 - \delta_t)(\pi^t, c_t) + \delta_t(\pi^{t+1}, c_{t+1}),$$

where  $\delta_t := p_1/p_t \in (0, 1]$ . As the action assigned to each contract  $t \notin C$  is the same as before, their contribution to the objective is unchanged. But the contracts in  $C$  now yield

$$\begin{aligned} \sum_{t \in C} p_t \hat{\pi}^t (y - w^t) &= \sum_{t \in C} p_t [(1 - \delta_t) \pi^t + \delta_t \pi^{t+1}] (y - w^t) \\ &= \sum_{t \in C} [(p_t - p_1) \pi^t + p_1 \pi^{t+1}] (y - w^t) \\ &= \sum_{t \in C} p_t \pi^t (y - w^t) + p_1 \sum_{t \in C} (\pi^t - \pi^{t+1}) (w^t - y) \\ &= \sum_{t \in C} p_t \pi^t (y - w^t) - p_1 m \varepsilon, \end{aligned}$$

where the last equality follows by (B.4) and the fact that  $\sum_{t \in C} (\pi^t - \pi^{t+1}) y = 0$  because the sum is cyclical. Therefore, the new solution  $\{(\hat{\pi}^t, \hat{c}_t)\}_{t \in T}$  is a strict improvement over the supposed optimal solution  $\{(\pi^t, c_t)\}_{t \in T}$ , provided we show that it is feasible.

To show feasibility, we note first that each  $(\hat{\pi}^t, \hat{c}_t)$  satisfies the probability and non-negativity constraints (3.13) and (3.14). For  $t \notin C$  this is trivial as  $(\hat{\pi}^t, \hat{c}_t) = (\pi^t, c_t)$ . For  $t \in C$  this follows because  $(\hat{\pi}^t, \hat{c}_t)$  is by definition a convex combination of two actions, each of which satisfies (3.13) and (3.14).

Note then that, given any contract  $w^t$ , the agent's payoff from the new action  $(\hat{\pi}^t, \hat{c}_t)$  is weakly higher than from the old action  $(\pi^t, c_t)$  by construction:

$$\hat{\pi}^t w^t - \hat{c}_t = \begin{cases} \pi^t w^t - c_t + \delta_t (\pi^{t+1} w^t - c_{t+1} - \pi^t w^t + c_t) = \pi^t w^t - c_t + \delta_t \varepsilon & \text{if } t \in C, \\ \pi^t w^t - c_t & \text{if } t \notin C, \end{cases}$$

where the first case follows by the definition of  $(\hat{\pi}^t, \hat{c}_t)$  and equation (B.3). This implies that the new solution satisfies the participation constraint (3.12) for all  $t$ , since

$$\hat{\pi}^t w^t - \hat{c}_t \geq \pi^t w^t - c_t \geq U^0(w^t).$$

It remains to verify the incentive compatibility constraint (B.1). Fix contracts  $t$  and  $s \neq t$ . Observe that by construction of  $(\hat{\pi}^s, \hat{c}_s)$ , we have  $(\hat{\pi}^s, \hat{c}_s) = (1 - \beta)(\pi^s, c_s) + \beta(\pi^{s+1}, c_{s+1})$  for

$\beta \in (0, 1]$  (if  $s \in C$ ) or for  $\beta = 0$  (if  $s \notin C$ ). Therefore,

$$\hat{\pi}^t w^t - \hat{c}_t \geq \pi^t w^t - c_t \geq (1 - \beta)(\pi^s w^t - c_s - \varepsilon) + \beta(\pi^{s+1} w^t - c_{s+1} - \varepsilon) = \hat{\pi}^s w^t - \hat{c}_s - \varepsilon,$$

where the second inequality follows because  $\{(\pi^t, c_t)\}_{t \in T}$  satisfies (B.1) by assumption.

We conclude that the new solution  $\{(\hat{\pi}^t, \hat{c}_t)\}_{t \in T}$  is feasible, a contradiction.  $\square$

To prove Lemma 2, for each  $n \in \mathbb{N}$ , let  $(\kappa_n^*, \lambda_n^*, \mu_n^*)$  be an optimal extreme point solution to the  $1/n$ -perturbed dual.<sup>16</sup> The feasible set is independent of the perturbation parameter and it contains only finitely many extreme points. Hence, one of them appears infinitely often as  $n \rightarrow \infty$ . Thus, by extracting a subsequence if necessary, we may assume that the sequence is constant. That is, there is an extreme point  $(\kappa^*, \lambda^*, \mu^*)$  such that  $(\kappa_n^*, \lambda_n^*, \mu_n^*) = (\kappa^*, \lambda^*, \mu^*)$  for all  $n$ . By Lemma B.1, the associated graph  $G(\kappa^*)$  is acyclic.

We claim that  $(\kappa^*, \lambda^*, \mu^*)$  is optimal in the 0-perturbed dual (3.15–3.18). To see this, note that the perturbed dual objective  $f(\kappa, \lambda, \mu, \varepsilon) := \sum_{t \in T} (\lambda_t U(w^t) + \mu_t) - \varepsilon \sum_{t \in T} \sum_{s \neq t} \kappa_{st}$  is jointly continuous in  $(\kappa, \lambda, \mu, \varepsilon)$ . Therefore, given any feasible solution  $(\kappa, \lambda, \mu)$ , we have

$$f(\kappa, \lambda, \mu, 0) = \lim_n f(\kappa, \lambda, \mu, 1/n) \leq \lim_n f(\kappa^*, \lambda^*, \mu^*, 1/n) = f(\kappa^*, \lambda^*, \mu^*, 0),$$

where the inequality is because  $(\kappa^*, \lambda^*, \mu^*)$  is optimal for all  $n$  by construction. We conclude that  $(\kappa^*, \lambda^*, \mu^*)$  is an optimal solution to (3.15–3.18) such that  $G(\kappa^*)$  is acyclic.

## B.2 Proof of Lemma 4

We start with a standard constraint simplification result:

**Lemma B.2.** *For any  $e(\cdot)$ , there exists  $c(\cdot)$  such that  $(e(\cdot), c(\cdot))$  is feasible in (5.4–5.6) if and only if  $e(\cdot)$  is feasible in (5.8–5.10).*

*Proof.* Let  $(e(\cdot), c(\cdot))$  be feasible in (5.4–5.6). Write  $U(\alpha) := \alpha e(\alpha) - c(\alpha)$  for the agent's indirect utility given contract  $\alpha$ . Because  $(e(\cdot), c(\cdot))$  satisfies (5.4),  $e(\cdot)$  satisfies (5.8) and the envelope theorem (Myerson, 1981; Milgrom and Segal, 2002) implies that

$$U(\alpha) = U(0) + \int_0^\alpha e(t) dt \quad \forall \alpha \in [0, 1]. \quad (\text{B.5})$$

---

<sup>16</sup>To see that the feasible set of the  $\varepsilon$ -perturbed dual has an extreme point (and thus an extreme point that is optimal, because the optimal value is finite) even though there are free variables, let  $\mu'_t := \min_{i \in I} p_t(y_i - w_i^t)$  for all  $t \in T$ . Then  $(\kappa, \lambda, \mu) = (0, 0, \mu')$  is an extreme point.

The feasibility constraints in (5.6) imply

$$U(0) \leq 0, \quad e(0) \geq 0, \quad e(1) \leq y_n,$$

which clearly imply (5.10). (Here  $U(0) \leq 0$  follows because  $U(0) = -c(0) \leq 0$ .) Finally, (B.5) and the participation constraints in (5.5) imply that

$$U(0) + \int_0^\alpha e(t)dt \geq U^0(\alpha) \quad \forall \alpha \in [0, 1],$$

which implies (5.9), because  $U(0) = -c(0) \leq 0$ . Thus  $(e(\cdot), c(\cdot))$  is feasible in (5.8–5.10).

The converse is shown by taking  $e(\cdot)$  feasible in (5.8–5.10) and constructing  $c(\cdot)$  via (B.5) as  $c(\alpha) = \alpha e(\alpha) - \int_0^\alpha e(t)dt$  (with  $c(0) = 0$ ). The details are standard and hence omitted.  $\square$

It remains to show that  $V(p) = L(p)$ . The inequality  $V(p) \geq L(p)$  follows, because any  $A \supseteq A^0$  gives rise to a feasible solution to (5.3–5.6), and hence  $V(p, A) \geq L(p)$ .

To show that  $V(p) \leq L(p)$ , we fix an optimal solution  $(e^*(\cdot), c^*(\cdot))$  to (5.3–5.6) and construct a feasible technology  $A$  such that  $V(p, A) = L(p)$ . The obvious candidate for  $A$  is the set  $A' := \{(e^*(\alpha), c^*(\alpha)) : \alpha \in [0, 1]\} \cup A^0$ , where each  $(e^*(\alpha), c^*(\alpha))$  is the equivalence class of all actions with expected output  $e^*(\alpha)$  and cost  $c^*(\alpha)$ . However,  $A'$  need not be compact if  $(e^*(\cdot), c^*(\cdot))$  is not continuous. Therefore, we take  $A$  to be the closure of  $A'$ . More specifically, the proof has three steps: 1) We show that  $e^*(\cdot)$  can be taken to be lower semi-continuous. 2) We then define  $A$  as the closure of  $A'$  and verify that this gives a feasible technology where  $(e^*(\alpha), c^*(\alpha))$  remains a best response to  $\alpha$  for all  $\alpha \in [0, 1]$ . Finally, 3) we verify that each  $(e^*(\alpha), c^*(\alpha))$  satisfies the adversarial tie-breaking assumption.

Let  $B(\alpha, \tilde{A}) := \arg \max_{(e,c) \in \tilde{A}} \alpha e - c$  be the set of (equivalence classes of) actions that are best responses to contract  $\alpha$  given any technology  $\tilde{A}$ . We say that a solution  $(e^*(\cdot), c^*(\cdot))$  to (5.3–5.6) satisfies adversarial tie-breaking at the bottom if  $e^*(0) = \min\{e : (e, c) \in B(0, A')\}$ .

**Lemma B.3.** *The minimization problem (5.3–5.6) has an optimal solution  $(e^*(\cdot), c^*(\cdot))$  such that  $e^*(\cdot)$  is lower semi-continuous and satisfies adversarial tie-breaking at the bottom.*

*Proof.* Consider first the equivalent problem (5.7–5.10). Let  $e(\cdot)$  be a minimizer and denote by  $D \subset [0, 1]$  the set of points at which  $e(\cdot)$  is discontinuous. Because  $e(\cdot)$  is nondecreasing,  $D$  has at most countably many elements.

Define  $e^*(\cdot) : [0, 1] \rightarrow [0, y_n]$  by letting  $e^*(\alpha) = e(\alpha)$  for all  $\alpha \notin D$  and  $e^*(\alpha) = \lim_{\alpha' \uparrow \alpha} e(\alpha')$  for all  $\alpha \in D$ . By construction,  $e^*(\cdot)$  is nondecreasing and lower semi-continuous, and it satisfies (5.10). Moreover, because  $e(\cdot)$  satisfies (5.9) and  $D$  is countable,  $e^*(\cdot)$  satisfies (5.9) as well. Finally, because we have  $e^*(\alpha) \leq e(\alpha)$  for all  $\alpha \in [0, 1]$  by construction, the

allocation rule  $e^*(\cdot)$  yields a weakly lower profit than the minimizer  $e(\cdot)$ . Therefore,  $e^*(\cdot)$  is a lower semi-continuous minimizer to problem (5.7–5.10).

By Lemma B.2, there exists  $c^*(\cdot)$  such that  $(e^*(\cdot), c^*(\cdot))$  is an optimal solution to (5.3–5.6). Because  $e^*(\cdot)$  is nondecreasing, it can violate adversarial tie-breaking at the bottom only if there exists a known action  $(e_0, c_0) \in A^0$  such that  $(e_0, c_0) \in B(0, A')$  and  $e_0 < e^*(0)$ . If so, it is clearly feasible to replace  $(e^*(0), c^*(0))$  with  $(e_0, c_0)$  to obtain a new optimal solution that satisfies adversarial tie-breaking at the bottom.  $\square$

**Lemma B.4.** *Let  $(e^*(\cdot), c^*(\cdot))$  be an optimal solution to problem (5.3)–(5.6). Let  $A$  be the closure of  $A' := \{(e^*(\alpha), c^*(\alpha)) : \alpha \in [0, 1]\} \cup A^0$ . Then  $A$  is compact and thus it is a feasible technology. Furthermore, we have  $(e^*(\alpha), c^*(\alpha)) \in B(\alpha, A)$  for all  $\alpha \in [0, 1]$ .*

*Proof.* We verify first that  $A$  is compact. Note that the space of output distributions,  $\Delta(Y)$ , is compact because  $Y$  is finite (and thus compact). Moreover, we have  $c^*(\alpha) \in [0, y_n - U^0(\alpha)]$  for all  $\alpha$  by (5.5) and (5.6), and thus the cost of each action in  $A'$  is bounded from above by  $\bar{c} := \max\{y_n - U^0(0), \max\{c : (\pi, c) \in A^0\}\}$ . Thus,  $A'$  is a subset of the compact set  $\Delta(Y) \times [0, \bar{c}]$ , and hence its closure  $A$  is compact.

To prove the second claim, we note that, given any  $\alpha \in [0, 1]$ , the agent's payoff  $\alpha(\pi y) - c$  is a continuous function of  $(\pi, c)$  that attains its maximum on  $A'$  by construction, and therefore the maximum cannot increase by taking the closure of  $A'$ .  $\square$

**Lemma B.5.** *Let  $(e^*(\cdot), c^*(\cdot))$  be an optimal solution to problem (5.3–5.6) such that  $e^*(\cdot)$  is lower semi-continuous and satisfies adversarial tie-breaking at the bottom. Define the feasible technology  $A$  as in Lemma B.4. Then for all  $\alpha \in [0, 1]$ , the action  $(e^*(\alpha), c^*(\alpha))$  yields the lowest profit to the principal across all best responses to contract  $\alpha$  given technology  $A$ , i.e.,*

$$(1 - \alpha)e^*(\alpha) = \min_{(e, c) \in B(\alpha, A)} (1 - \alpha)e \quad \forall \alpha \in [0, 1],$$

and hence  $V(p, A) = L(p)$ .

*Proof.* Suppose toward contradiction that for some  $\hat{\alpha} \in [0, 1]$ , the best response set  $B(\hat{\alpha}, A)$  contains an action  $(\hat{e}, \hat{c})$  such that  $(1 - \hat{\alpha})\hat{e} < (1 - \hat{\alpha})e^*(\hat{\alpha})$ .

Suppose first that  $\hat{\alpha} = 0$  so that  $e^*(0) > \hat{e}$ . Because  $e^*(\cdot)$  is nondecreasing, this implies  $(\hat{e}, \hat{c}) \in A^0$ . But then  $e^*(\cdot)$  violates adversarial tie-breaking at the bottom, a contradiction.

Suppose then that  $\hat{\alpha} > 0$ . Because the agent's payoff  $\alpha e - c$  has strictly increasing differences in  $(\alpha, e)$ , the Monotone Selection Theorem implies that the selections  $(e^*(\alpha), c^*(\alpha)) \in B(\alpha, A)$  and  $(\hat{e}, \hat{c}) \in B(\hat{\alpha}, A)$  satisfy  $e^*(\alpha) \leq \hat{e}$  for all  $\alpha < \hat{\alpha}$ . Thus,  $\lim_{\alpha \uparrow \hat{\alpha}} e^*(\alpha) \leq \hat{e} < e^*(\hat{\alpha})$ , and hence  $e^*(\cdot)$  is not lower semi-continuous at  $\hat{\alpha}$ , a contradiction.  $\square$

To conclude the proof, by Lemma B.3 we can take an optimal solution to (5.3–5.6) where  $e^*(\cdot)$  is lower semi-continuous and satisfies adversarial tie-breaking at the bottom. Lemma B.5 then gives a technology  $A$  such that  $V(p, A) = L(p)$ , which implies  $V(p) \leq L(p)$ .

### B.3 Proof of Corollary 3

Denote the maximands in the optimal guarantee (5.1) and Carroll's (2015) optimal deterministic guarantee, reproduced for adversarial tie-breaking in (A.2) in Appendix A, by

$$f(\alpha) := \frac{U^0(\alpha)}{-\ln(1-\alpha)} \quad \text{and} \quad g(\alpha) := \frac{1-\alpha}{\alpha} U^0(\alpha).$$

Recalling that these quotients are extended to all of  $[0, 1]$  by left- and right-continuity, it is straightforward to verify that this defines continuous functions  $f, g : [0, 1] \rightarrow \mathbb{R} \cup \{-\infty\}$  such that  $f(0) = g(0)$  and  $f(1) = g(1) = 0$ . (The first equality follows because  $f(\alpha)/g(\alpha) \rightarrow 1$  as  $\alpha \rightarrow 0$  by l'Hospital's rule.) Moreover, because  $(1-\alpha)/\alpha < -1/\ln(1-\alpha)$  for all  $\alpha \in (0, 1)$ , we have

$$g(\alpha) > 0 \implies f(\alpha) > g(\alpha) \quad \forall \alpha \in (0, 1). \quad (\text{B.6})$$

Suppose now that the inequality in Corollary 3 is not satisfied. Then there is no gain from randomization, because  $\max_{\alpha} f(\alpha) = f(0) = g(0) \leq \max_{\alpha} g(\alpha)$ .

In the other direction, suppose the inequality in Corollary 3 is satisfied. Let  $\alpha^* > 0$  attain the maximum on the right-hand side (and hence that of  $f$ ) and let  $\alpha_D^*$  maximize  $g$ . Because the known technology contains a surplus-generating action, we have  $\alpha_D^* < 1$  and  $g(\alpha^*) > 0$ . If  $\alpha_D^* = 0$ , then  $g(\alpha^*) = g(0) = f(0) < f(\alpha^*)$ . If instead  $\alpha_D^* > 0$ , then (B.6) implies that  $f(\alpha^*) \geq f(\alpha_D^*) > g(\alpha_D^*)$ . We conclude that either way,  $\max_{\alpha} f(\alpha) > \max_{\alpha} g(\alpha)$ , i.e., the gain from randomization is positive.

### B.4 Proof of Proposition 4

1. The optimal guarantee from Proposition 3 becomes

$$\max_{\alpha \in [0, 1]} \frac{\alpha - c}{-\ln(1-\alpha)}.$$

The necessary first-order condition to this problem is

$$-\ln(1-\alpha) - (1-\alpha)^{-1}(\alpha - c) = 0. \quad (\text{B.7})$$

It is readily verified to have a unique solution,  $\alpha^*(c)$ , which is a continuous increasing function of  $c$ , with  $\lim_{c \rightarrow 0} \alpha^*(c) = 0$  and  $\lim_{c \rightarrow 1} \alpha^*(c) = 1$ . Moreover, the left-hand side crosses zero once from above. To show that  $\alpha^*(c)$  is strictly larger than the slope of the optimal deterministic contract,  $\alpha_D(c) = \sqrt{c}$ , it thus suffices to show that the left-hand side of (B.7) is strictly larger than zero when evaluated at  $\sqrt{c}$ :

$$-\ln(1 - \sqrt{c}) - \frac{\sqrt{c} - c}{1 - \sqrt{c}} = -\ln(1 - \sqrt{c}) - \sqrt{c} > 0.$$

It remains to establish that  $\alpha^*(c)$  is strictly increasing and concave. To show this, observe that the implicit function theorem implies that the derivative of  $\alpha^*(c)$  on  $(0, 1)$  is  $(\alpha^*)'(c) = -1/\ln(1 - \alpha^*(c)) > 0$ . Moreover, its second derivative is given by  $(\alpha^*)''(c) = -(\alpha^*)'(c)/((1 - \alpha^*(c))(\ln(1 - \alpha^*(c)))^2) < 0$ .

2. Using (B.7), it is easy to show that the optimal guarantee is  $\bar{v}(c) = 1 - \alpha^*(c)$ . Because  $(\alpha^*)'(c) > 0$  and  $(\alpha^*)''(c) < 0$ ,  $\bar{v}(c)$  is decreasing and convex. Moreover,  $\lim_{c \rightarrow 0} \alpha^*(c) = 0$  and  $\lim_{c \rightarrow 1} \alpha^*(c) = 1$  implies  $\lim_{c \rightarrow 0} \bar{v}(c) = 1$  and  $\lim_{c \rightarrow 1} \bar{v}(c) = 0$ .
3. A straightforward calculation using (A.2) gives the optimal deterministic guarantee  $(1 - \sqrt{c})^2$ . Hence, the gain from randomization is

$$\gamma(c) = \frac{1 - \alpha^*(c)}{(1 - \sqrt{c})^2}.$$

From  $\lim_{c \rightarrow 0} \alpha^*(c) = 0$ , it is immediate that  $\lim_{c \rightarrow 0} \gamma(c) = 1$ . From  $\lim_{c \rightarrow 1} \alpha^*(c) = 1$  and twice applying l'Hôpital's rule, we have  $\lim_{c \rightarrow 1} \gamma(c) = \infty$ .

## Appendix C: The Gilboa–Schmeidler axiomatization

We show that the principal's maxmin criterion can be interpreted as uncertainty aversion in the sense of Gilboa and Schmeidler (1989, GS henceforth).

Recall the Anscombe–Aumann setting considered by GS: There is a set  $S$  of *states*, endowed with an algebra of events  $\Sigma$ , and a set  $X$  consisting of all finitely supported distributions over some underlying set  $\Omega$  of *outcomes*.<sup>17</sup> The elements of  $X$  are referred to as (*roulette*) *lotteries*. An *act* (or a horse lottery) is a  $\Sigma$ -measurable function  $f : S \rightarrow X$ . We identify each lottery  $x \in X$  with the constant act equal to  $x$ . Convex combinations of acts are evaluated pointwise: given acts  $f, g$  and a number  $\eta \in [0, 1]$ , the act  $h := \eta f + (1 - \eta)g$  is defined by  $h(s) := \eta f(s) + (1 - \eta)g(s)$  for all  $s \in S$ . In other words, a randomization

---

<sup>17</sup>We adopt notation different from that used by GS in order to avoid conflict with the rest of the paper.

between acts is simply another act. Let  $\mathcal{F}_0$  denote the set of all simple acts (i.e., acts whose range is finite), and let  $\mathcal{F}$  be any convex set of acts that contains  $\mathcal{F}_0$ .

Let  $\succsim$  be a binary relation on  $\mathcal{F}$ . The GS axioms on  $\succsim$  are the following:

- A1. (*Weak order*)  $\succsim$  is complete and transitive.
- A2. (*Certainty independence*) For all  $f, g \in \mathcal{F}$ ,  $x \in X$ , and  $\eta \in (0, 1)$ ,  $f \succ g$  if and only if  $\eta f + (1 - \eta)x \succ \eta g + (1 - \eta)x$ .
- A3. (*Continuity*) For all  $f, g, h \in \mathcal{F}$ , if  $f \succ g$  and  $g \succ h$ , then there exist  $\eta, \zeta \in (0, 1)$  such that  $\eta f + (1 - \eta)h \succ g$  and  $g \succ \zeta f + (1 - \zeta)h$ .
- A4. (*Monotonicity*) For all  $f, g \in \mathcal{F}$ , if  $f(s) \succsim g(s)$  for all  $s \in S$ , then  $f \succsim g$ .
- A5. (*Uncertainty aversion*) For all  $f, g \in \mathcal{F}$  and  $\eta \in (0, 1)$ , if  $f \sim g$ , then  $\eta f + (1 - \eta)g \succsim f$ .
- A6. (*Non-degeneracy*) It is not the case that  $f \succsim g$  for all  $f, g \in \mathcal{F}$ .

Given a weak order  $\succsim$  on  $\mathcal{F}$ , an act  $f \in \mathcal{F}$  is  $\succsim$ -*bounded* if there are lotteries  $x', x'' \in X$  such that  $x' \succsim f(s) \succsim x''$  for all  $s \in S$ . Let  $\mathcal{F}(\succsim)$  be the set of all  $\succsim$ -bounded acts. Clearly  $\mathcal{F}(\succsim)$  contains  $\mathcal{F}_0$ ; it is also convex if  $\succsim$  satisfies A1–A6.

The following representation theorem can be deduced from GS's Theorem 1 and the remark immediately after their Proposition 4.1.

**Proposition C.1** (Gilboa and Schmeidler, 1989). *Let  $\succsim$  be a binary relation on  $\mathcal{F}$ . Then the following conditions are equivalent:*

1. *The restriction of  $\succsim$  to  $\mathcal{F}(\succsim)$  satisfies axioms A1–A6.*
2. *There exist a non-constant affine function  $u : X \rightarrow \mathbb{R}$  and a nonempty, closed, convex set  $\mathcal{C}$  of finitely additive probability measures on  $(S, \Sigma)$  such that, for all  $f, g \in \mathcal{F}(\succsim)$ ,*

$$f \succsim g \quad \text{if and only if} \quad \min_{\nu \in \mathcal{C}} \int_S u(f(s)) d\nu(s) \geq \min_{\nu \in \mathcal{C}} \int_S u(g(s)) d\nu(s).$$

The contracting problem introduced in Section 2 can be viewed as a special case of this representation as follows. Let  $S = \{A : A \supseteq A^0\}$  with  $\Sigma = 2^S$ , so that a state is a technology containing the known technology. Take the outcomes to be profit levels,  $\Omega = \mathbb{R}$ , so that  $X = \Delta(\mathbb{R})$  is the set of finitely supported profit lotteries. As the principal is risk-neutral with respect to these objective lotteries, her preferences over  $X$  are represented by the linear (and thus affine) utility

$$u(x) = \int_{\mathbb{R}} z dx(z).$$

Every deterministic contract  $w \in \mathbb{R}_+^n$  induces an act  $f_w : S \rightarrow X$ , where for each technology  $A \in S$ ,  $f_w(A) \in X$  is the profit lottery generated by  $w$  when the agent chooses an adversarial best response from  $A$ . (Since the output set  $Y$  is finite,  $f_w(A)$  has finite support.) By (2.3) we have

$$u(f_w(A)) = V(w|A).$$

A finitely supported random contract  $p \in \Delta(\mathbb{R}_+^n)$  then induces an act  $f_p$  that is a convex combination of the acts  $\{f_w : w \in \text{supp}(p)\}$ . That is, for all  $A \in S$ ,

$$f_p(A) = \sum_{w \in \text{supp}(p)} p(w) f_w(A),$$

and by linearity of  $u$ ,

$$u(f_p(A)) = \sum_{w \in \text{supp}(p)} p(w) u(f_w(A)) = \sum_{w \in \text{supp}(p)} p(w) V(w|A) = V(p|A),$$

where the last equality is by (2.5). Note that  $u \circ f_p$  is bounded: given any technology  $A \in S$ , we have  $\min_{w \in \text{supp}(p)} \min_{i \in I} (y_i - w_i) \leq u(f_p(A)) \leq \max_{w \in \text{supp}(p)} \max_{i \in I} (y_i - w_i)$ . It follows that each  $f_p$  is  $\succsim$ -bounded for  $\succsim$  induced by (2.6).

Let  $\mathcal{D}$  be the set of Dirac measures on  $S$  and  $\text{co } \mathcal{D}$  its convex hull. Denote by  $\mathcal{C}$  the closure of  $\text{co } \mathcal{D}$  in the weak-\* topology on finitely additive probability measures on  $(S, \Sigma)$ . By the Banach–Alaoglu theorem,  $\mathcal{C}$  is weak-\* compact. Moreover, the map  $\nu \mapsto \int u \circ f_p d\nu$  is weak-\* continuous since  $u \circ f_p$  is bounded. Thus, it attains a minimum on  $\mathcal{C}$ , which satisfies

$$\min_{\nu \in \mathcal{C}} \int_S u(f_p(A)) d\nu(A) = \inf_{\nu \in \text{co } \mathcal{D}} \int_S u(f_p(A)) d\nu(A) = \inf_{\nu \in \mathcal{D}} \int_S u(f_p(A)) d\nu(A) = \inf_{A \in S} u(f_p(A)).$$

Consequently, the principal's guarantee defined in (2.6) can be written as

$$V(p) = \inf_{A \in S} V(p|A) = \min_{\nu \in \mathcal{C}} \int_S u(f_p(A)) d\nu(A),$$

i.e., it is a GS representation in the sense of Proposition C.1.<sup>18</sup>

---

<sup>18</sup>To minimize technicalities, we have here restricted attention to finitely supported random contracts. General random contracts can be handled by replacing  $X = \Delta(\mathbb{R})$  by a barycentric consequence space such as  $\bar{\Delta}_1(\mathbb{R})$ , the set of Borel probability measures on profits with finite first moment, and defining  $f_p(A)$  as the barycenter  $f_p(A) = \int f_w(A) p(dw)$ . Under appropriate measurability and integrability conditions (e.g., if  $w \mapsto f_w(A)$  is Borel and  $\int \|w\| p(dw) < \infty$ ), this barycenter is well defined and satisfies

$$u(f_p(A)) = \int u(f_w(A)) p(dw) = \int V(w|A) p(dw).$$

The above GS representation can be extended to this setting.

## References

- ANSCOMBE, F. J., AND R. J. AUMANN (1963): “A Definition of Subjective Probability,” *The Annals of Mathematical Statistics*, 34(1), 199–205.
- ANTIC, N. (2021): “Contracting with Unknown Technologies,” Discussion paper, Northwestern University.
- ANTIC, N., AND G. GEORGIADIS (2024): “Robust Contracts: A Revealed Preference Approach,” *The Review of Economics and Statistics*, forthcoming.
- BROOKS, B., AND S. DU (2021): “Optimal Auction Design with Common Values: An Informationally Robust Approach,” *Econometrica*, 89(3), 1313–1360.
- BROOKS, B., AND S. DU (2024): “On the Structure of Informationally Robust Optimal Mechanisms,” *Econometrica*, 92, 1391–1438.
- BROOKS, B., AND S. DU (2025): “Dual Reductions and the First-Order Approach for Informationally Robust Mechanism Design,” Discussion paper, University of Chicago.
- BURKETT, J., AND M. ROSENTHAL (2024): “Data-Driven Contract Design,” *Journal of Economic Theory*, 221, 105585.
- CARRASCO, V., V. F. LUZ, N. KOS, M. MESSNER, P. MONTEIRO, AND H. MOREIRA (2018): “Optimal selling mechanisms under moment conditions,” *Journal of Economic Theory*, 177, 245–279.
- CARROLL, G. (2015): “Robustness and Linear Contracts,” *American Economic Review*, 105(2), 536–563.
- (2019): “Robustness in Mechanism Design and Contracting,” *Annual Review of Economics*, 11, 139–166.
- CARROLL, G., AND L. BOLTE (2023): “Robust contracting under double moral hazard,” *Theoretical Economics*, 18(4), 1623–1663.
- CHE, Y.-K., AND W. ZHONG (2024): “Robustly Optimal Mechanisms for Selling Multiple Goods,” *The Review of Economic Studies*, forthcoming.
- DAI, T., AND J. TOIKKA (2022): “Robust Incentives for Teams,” *Econometrica*, 90(4), 1583–1613.

- DEB, R., AND A.-K. ROESLER (2024): “Multi-Dimensional Screening: Buyer-Optimal Learning and Informational Robustness,” *The Review of Economic Studies*, 91, 2744–2770.
- ELLSBERG, D. (1961): “Risk, Ambiguity, and the Savage Axioms,” *The Quarterly Journal of Economics*, 75(4), 643–669.
- GILBOA, I., AND D. SCHMEIDLER (1989): “Maxmin expected utility with non-unique prior,” *Journal of mathematical economics*, 18(2), 141–153.
- HANANY, E., AND P. KLIBANOFF (2025): “Robust Contracting and Voluntary Disclosure,” Discussion paper, Northwestern University.
- HOLMSTRÖM, B., AND P. MILGROM (1987): “Aggregation and Linearity in the Provision of Intertemporal Incentives,” *Econometrica*, 55(2), 303–328.
- KAMBHAMPATI, A. (2023): “Randomization Is Optimal in the Robust Principal-Agent Problem,” *Journal of Economic Theory*, 207, 105585.
- (2024): “Robust Performance Evaluation of Independent Agents,” *Theoretical Economics*, 19(3).
- KAMBHAMPATI, A., J. TOIKKA, AND R. VOHRA (2024): “Randomization and the Robustness of Linear Contracts,” in *North American Winter Meeting of the Econometric Society, San Antonio, Texas*.
- KE, S., AND Q. ZHANG (2020): “Randomization and Ambiguity Aversion,” *Econometrica*, 88(3), 1159–1195.
- LAI, H.-C., AND S.-Y. WU (1992): “On Linear Semi-Infinite Programming Problems: An Algorithm,” *Numerical Functional Analysis and Optimization*, 13(3–4), 287–304.
- MARKU, K., S. OCAMPO DIAZ, AND J.-B. TONDJI (2024): “Robust Contracts in Common Agency,” *The Rand Journal of Economics*, 55(2), 199–229.
- MILGROM, P., AND I. SEGAL (2002): “Envelope theorems for arbitrary choice sets,” *Econometrica*, 70(2), 583–601.
- MYERSON, R. B. (1981): “Optimal Auction Design,” *Mathematics of Operations Research*, 6(1), 58–73.
- NÖLDEKE, G., AND L. SAMUELSON (2018): “The implementation duality,” *Econometrica*, 86(4), 1283–1324.

- OLSZEWSKI, W. (2025): “On evaluating contracts by their performance in the worst cases,” Discussion paper, Northwestern University.
- PENG, B., AND Z. G. TANG (2024): “Optimal Robust Contract Design,” in *ACM Conference on Economics and Computation (EC’24)*.
- RAIFFA, H. (1961): “Risk, Ambiguity, and the Savage Axioms: Comment,” *The Quarterly Journal of Economics*, 75(4), 690–694.
- ROCHET, J.-C. (2024): “Multidimensional Screening After 37 Years,” *Journal of Mathematical Economics*, 113, 103010.
- ROSENTHAL, M. (2023): “Robust Incentives for Risk,” *Journal of Mathematical Economics*, 109, 102893.
- (2025): “Simple Incentives and Diverse Beliefs,” *Theoretical Economics*, forthcoming.
- SAITO, K. (2015): “Preferences for Flexibility and Randomization Under Uncertainty,” *American Economic Review*, 105(3), 1246–71.
- VAIRO, M. (2025): “Robustly Optimal Income Taxation,” Discussion paper, Princeton University.
- WALTON, D., AND G. CARROLL (2022): “A General Framework for Robust Contracting Models,” *Econometrica*, 90(5), 2129–2159.