

A general solution method for moral hazard problems

RONGZHU KE

Department of Economics, Hong Kong Baptist University

CHRISTOPHER THOMAS RYAN

Booth School of Business, University of Chicago

Principal–agent models are pervasive in theoretical and applied economics, but their analysis has largely been limited to the “first-order approach” (FOA), where incentive compatibility is replaced by a first-order condition. This paper presents a new approach to solving a wide class of principal–agent problems that satisfy the monotone likelihood ratio property but may fail to meet the requirements of the FOA. Our approach solves the problem via tackling a max–min–max formulation over agent actions, alternate best responses by the agent, and contracts.

KEYWORDS. Principal–agent, moral hazard, solution method.

JEL CLASSIFICATION. D82, D86.

1. INTRODUCTION

Moral hazard principal–agent problems are well studied, but unresolved technical difficulties persist. An essential difficulty is finding a tractable method to deal with the incentive compatibility (IC) constraints that capture the strategic behavior of the agent. Incentive compatibility is challenging for at least two reasons. First, when the agent’s action space is continuous, there are, in principle, infinitely many IC constraints. Second, these constraints turn the principal’s decision into an optimization problem over a potentially nonconvex set. Much attention has been given to finding structure in special cases that overcome these issues. The *first-order approach* (FOA), where the IC constraints are replaced by the first-order condition of the agent’s problem (Rogerson 1985, Jewitt 1988), applies when the agent’s objective function is concave in the agent’s action. Previous studies have proposed various sufficient conditions for the FOA to be valid (see, e.g., Rogerson 1985, Jewitt 1988, Sinclair-Desgagné 1994, Conlon 2009, Kirkegaard 2017b, Jung and Kim 2015). Nonetheless, there remain natural settings where the FOA is invalid (see, for instance, Example 5 below).

When the FOA is invalid, more elaborate methods have been proposed.¹ Grossman and Hart (1983) explore settings where there are finitely many output scenarios and reduce incentive compatibility to a finite number of constraints. However, their method

Rongzhu Ke: rongzhuke@hkbu.edu.hk

Christopher Thomas Ryan: chris.ryan@chicagobooth.edu

We are indebted to the comments of two anonymous referees whose feedback greatly improved this paper.

¹Several authors have analyzed the moral hazard problem without the FOA in specialized settings. This includes Innes (1990), Wang and Hu (2016), and Kirkegaard (2017a).

does not apply when the agent's output takes on infinitely many values. An alternate approach is due to [Mirrlees \(1999\)](#) (which originally appeared in 1975) and is refined in [Mirrlees \(1986\)](#) and [Araujo and Moreira \(2001\)](#). This method overcomes the limitations of the FOA by reintroducing a subset of IC constraints, in addition to the first-order condition, to eliminate alternate best responses. These reintroduced constraints—called *no-jump constraints*—isolate attention to contract–action pairs that are incentive compatible. The main difficulty in Mirrlees's approach is in producing the required no-jump constraints. There is a potential to reintroduce many—if not infinitely many—no-jump constraints. Moreover, a general method for generating these constraints is not known and brute force enumeration is often difficult. [Araujo and Moreira \(2001\)](#) use second-order information to refine the search, but the essential difficulties remain.

The procedure described in this paper systematically builds on Mirrlees's approach. The problem of determining which no-jump constraints are needed is recast as a minimization problem that identifies the hardest-to-satisfy no-jump constraint over the set of alternate best responses. This makes the original problem equivalent to an optimization problem that involves three sequential optimal decisions: maximizing over the contract, maximizing over the agent's action, and minimizing over alternate best responses to that chosen action. We then propose a tractable relaxation to this problem by changing the order of optimization to max-min-max, where the former maximization is over agent actions and the latter maximization is over contracts. The analytical benefits of this new order are clear. The map that describes which optimal contracts support a given action against deviation to a specific alternate best response has desirable topological properties, which are explored in [Section 3](#). We call this max-min-max relaxation the *sandwich* relaxation since the inner minimization is “sandwiched” between two outer maximizations.

The main technical work of the paper is to uncover when the sandwich relaxation is tight. This involves careful consideration of what utility can be guaranteed to the agent by an optimal contract. In particular, if the individual rationality constraint is *not* binding, a family of sandwich relaxations indexed by lower bounds on agent utility that are larger than the reservation utility must be examined so as to find a relaxation that is tight. Constructing the appropriate bound and guaranteeing that the resulting relaxation is tight comprise a main focus of our development. Our development assumes monotonicity conditions on the output distribution; namely, the monotone likelihood ratio property (MLRP) that is defined carefully in the main body of the text.

It should be noted that the MLRP assumption is common in the usual discussion of the FOA. However, it is also well known that the MLRP is *insufficient* to guarantee the validity of the FOA ([Grossman and Hart 1983](#), [Rogerson 1985](#), [Jewitt 1988](#), [Conlon 2009](#)). We illustrate scenarios where the sandwich approach is valid (that is, the sandwich relaxation is tight) but the FOA is invalid. This is carefully discussed in [Section 5](#) where it is established that the sandwich approach ensures a stationarity condition for a worst-case alternate best response that is stronger than the stationarity condition in the FOA. This is due to the inner minimization over alternate best responses in the sandwich approach that is absent from the FOA. However, when the FOA is valid,

the sandwich approach is also valid and both approaches result in the same optimal contract.

Finally, we comment here on some similarities with a related paper written by the authors. In [Ke and Ryan \(2018\)](#), we consider a similar problem setting with similar assumptions. The main focus of that paper is to establish an important structural result, namely to recover a monotonicity result for optimal contracts under MLRP that holds even when the FOA is invalid. To that end, that paper takes the approach of [Grossman and Hart \(1983\)](#) of taking the agent's action as given and finds structure on those optimal contracts that implements the given action. Consequently, [Ke and Ryan \(2018\)](#) do not provide a general solution procedure for moral hazard problems, and instead focus on establishing structural properties of optimal contracts without explicitly constructing such policies. By contrast, the current paper is focused on the full problem that allows the agent's action to respond optimally to an offered contract. Of course, this adds significant complication to the analysis, hence the need for another paper. Indeed, consider the classical example of [Mirrlees \(1999\)](#) that initiated discussion of the failure of the FOA. If a tight reservation utility and best response are known, a first-order condition is easily shown to suffice. The failure of the FOA is precisely its inability to identify a target action of the follower. See also our [Example 1](#) and [Proposition 5](#) below for a related discussion.

A more subtle technical challenge here concerns questions of existence. The inner minimization in the sandwich problem need not be attained, an issue that is precluded from the analysis of [Ke and Ryan \(2018\)](#). There, a target best response a^* is specified and an assumption is made so that an alternate and distinct best response $\hat{a}^*$ exists. This assumption essentially rules out the validity of the FOA. In other words, the analysis of [Ke and Ryan \(2018\)](#) does not apply to many problems where the FOA is known to be valid. This is not a concern in that paper, since the goal there is to devise the structure of optimal contracts, particularly monotonicity properties, which are already known in the setting where the FOA is valid ([Rogerson 1985](#)). In contrast, the goal of this paper is to develop a general procedure for solving moral hazard problems that satisfy the MLRP, and thus should incorporate cases where the FOA additionally holds. [Section 5](#) provides more detail on how this existence issue is connected to the FOA.

Although there are similarities in the development of both papers (the current paper and [Ke and Ryan 2018](#)), they can largely be read independently. [Ke and Ryan \(2018\)](#) does not reference the current paper, and there are only several references to [Ke and Ryan \(2018\)](#) here, all of which appear in the Technical Appendix.²

This paper is organized as follows. [Section 2](#) contains the model and reviews existing approaches to solve the principal–agent problem. [Section 3](#) describes the sandwich relaxation and gives sufficient conditions for the relaxation to be tight, given an appropriately chosen lower bound on agent utility. [Section 4](#) describes the methodology to construct such lower bounds. [Section 5](#) discusses existence and its connection the FOA. [Section 6](#) provides three examples that illustrate the mechanics of our procedure. We consider a simplified moral hazard example throughout the paper to illuminate the theory. Proofs of all technical results are contained in the [Appendix](#).

²We thank an anonymous referee for raising and shedding light on the question of existence during the review process of the paper. We also thank another anonymous referee for drawing attention to the similarities and distinctions between the current paper and [Ke and Ryan \(2018\)](#).

2. MODEL AND EXISTING APPROACHES

2.1 *Principal–agent model*

We study the classic moral hazard principal–agent problem with a single task and single-dimensional output. An agent chooses an action $a \in \mathbb{A}$ that is unobservable to a principal. This action influences the random outcome $X \in \mathcal{X}$ through the probability density function $f(x, a)$, where x is an outcome realization. The principal chooses a wage contract $w : \mathcal{X} \rightarrow [\underline{w}, \infty)$, where $\underline{w}$ is an exogenously given minimum wage. The value of output to the principal obeys the function $\pi : \mathcal{X} \rightarrow \mathbb{R}$.

Given an outcome realization $x \in \mathcal{X}$, the agent and principal derive the following utilities. The agent's utility under action a is separable in wage $w(x)$ and action cost $c(a)$. In particular, he derives utility $u(w(x)) - c(a)$, where $u : [\underline{w}, \infty) \rightarrow \mathbb{R}$ and $c : \mathbb{A} \rightarrow \mathbb{R}$. The principal's utility for outcome x is a function of the net value $\pi(x) - w(x)$ and is denoted $v(\pi(x) - w(x))$, where $v : \mathbb{R} \rightarrow \mathbb{R}$. The agent's expected utility is $U(w, a) = \int u(w(x))f(x, a) dx - c(a)$ and the principal's expected utility is $V(w, a) = \int v(\pi(x) - w(x))f(x, a) dx$. The agent has an outside option worth utility $\underline{U}$.

The principal faces the optimization problem³

$$\max_{w \geq \underline{w}, a \in \mathbb{A}} V(w, a) \quad (\text{P})$$

subject to the conditions

$$U(w, a) \geq \underline{U}, \quad (\text{IR})$$

$$U(w, a) - U(w, \hat{a}) \geq 0 \quad \text{for all } \hat{a} \in \mathbb{A}, \quad (\text{IC})$$

where (IR) is the individual rationality constraint that guarantees participation of the agent by furnishing at least the reservation utility $\underline{U}$ and (IC) are the incentive compatibility constraints that ensure the agent responds optimally.

ASSUMPTION 1. *The following statements hold:*

A1.1. The outcome set $\mathcal{X}$ is an interval in $\mathbb{R}$ and the action set is the bounded interval $\mathbb{A} \equiv [\underline{a}, \bar{a}]$.

A1.2. The outcome X is a continuous random variable, and $f(x, a)$ is continuous in x and twice continuously differentiable in $a \in \mathbb{A}$.

A1.3. For $a, a' \in \mathbb{A}$ with $a \neq a'$, there exists a positive measure subset of $\mathcal{X}$ such that $f(x, a) \neq f(x, a')$.

A1.4. The support of $f(\cdot, a)$ does not depend on a and, hence (without loss of generality), we assume the support is $\mathcal{X}$ for all a .

A1.5. Wage w is a measurable function on $\mathcal{X}$.

³The notation $w \geq \underline{w}$ is shorthand for expressing $w(x) \geq \underline{w}$ for almost all $x \in \mathcal{X}$.

A1.6. The value function π is increasing, continuous, and almost everywhere differentiable.

A1.7. The expected value $\int \pi(x)f(x, a) dx$ of output is bounded for all a .

A1.8. The agent's cost function c is increasing and continuously differentiable in a .

A1.9. The agent's utility function u is continuously differentiable, increasing, and strictly concave.

A1.10. The principal's utility function v is continuously differentiable, increasing, and concave.

The above assumptions are standard, so we do not spend time to justify them here.

ASSUMPTION 2. We also make the following additional technical assumptions:

A2.1. Either $\lim_{y \rightarrow \infty} u(y) = \infty$ or $\lim_{y \rightarrow -\infty} v(y) = -\infty$.

A2.2. The minimum wage $\underline{w}$, reservation utility $\underline{U}$, and least costly action $\underline{a}$ are such that $u(\underline{w}) - c(\underline{a}) < \underline{U}$.

The two conditions in this assumption are required in the proof of Lemma 3 that uses a Lagrangian duality method and ensures the existence of optimal dual solutions. Finally, to focus the scope of our paper, we make one additional assumption.

ASSUMPTION 3. There exists an optimal solution to (P). Moreover, we assume that the first-best contract is not optimal.

Existence is a challenging issue in its own right and is not the focus of this paper. We are interested in how to construct an optimal solution when one is known to exist. Several existing papers pay careful attention to the issue of existence. For instance, Kadan et al. (2017) provide weak sufficient conditions that guarantee the existence of an optimal solution. Moreover, we may assume that the first-best contract is not optimal without loss of interest, since finding a first-best contract is a well understood problem that not worthy of additional consideration.

We use the following terminology and notation. Let $a^{\text{BR}}(w)$ denote the set of actions that satisfy the (IC) constraint for a given contract w . That is, $a^{\text{BR}}(w) \equiv \arg \max_{a'} U(w, a')$. Let $\mathcal{F}$ denote the set of feasible solutions to (P). That is,

$$\mathcal{F} \equiv \{(w, a) : w \geq \underline{w}, a \in a^{\text{BR}}(w), U(w, a) \geq \underline{U}\}.$$

Given an action a , contract w is said to *implement* a if $(w, a) \in \mathcal{F}$. An action a is *implementable* if there exists a w that implements a . Let $\text{val}(\ast)$ denote the optimal value of the optimization problem $(\ast)$. In particular, $\text{val}(\text{P})$ denotes the optimal value of the original moral hazard problem (P). The single constraint in (IC) of the form

$$U(w, a) - U(w, \hat{a}) \geq 0 \quad (\text{NJ}(a, \hat{a}))$$

is called the *no-jump* constraint at $\hat{a}$ given a .

2.2 Existing approaches

We discuss the approaches to solve (P) that appear in the literature and their limitations. The standard-bearer is the FOA, which replaces (IC) with first-order conditions. Every implementable action a is an optimizer of the agent's problem and so satisfies necessary optimality conditions for that problem. In particular, a satisfies the first-order necessary condition

$$\begin{aligned} U_a(w, a) &= 0 & \text{if } a \in (\underline{a}, \bar{a}), & & U_a(w, a) \leq 0 & \text{if } a = \underline{a}, & \text{ and} \\ U_a(w, a) &\geq 0 & \text{if } a = \bar{a}, & \end{aligned} \quad (\text{FOC}(a))$$

where the subscripts denote partial derivatives. Replacing (IC) with (FOC(a)), problem (P) becomes

$$\max_{w \geq \underline{w}, a \in \mathbb{A}} \{V(w, a) : U(w, a) \geq \underline{U} \text{ and } (\text{FOC}(a))\}. \quad (\text{FOA})$$

When (FOA) and (P) have the same value (that is, $\text{val}(\text{P}) = \text{val}(\text{FOA})$) and the solution (w, a) to (FOA) has a implemented by w , we say the FOA is *valid*. Otherwise, the FOA is *invalid*.

Following Mirrlees (1999), we consider a special (very simplified) case of the moral hazard model that facilitates a geometric understanding of the technical issues involved. We return to this example at several points throughout the paper to ground our intuition. Section 6 has three additional examples of more general moral hazard problems that provide additional insights.

EXAMPLE 1. Suppose the principal chooses contract $z \in \mathbb{R}$ (following Mirrlees 1999, we use z to denote a single-dimensional contract instead of w) and the agent chooses an action $a \in [-2, 2]$ with reservation utility $\underline{U} = -2$. There is no lower bound on z . The principal obtains utility $v(z, a) = za - 2a^2$ and the agent receives benefit $-za$, minus action cost $c(a) = (a^2 - 1)^2$, with total agent utility

$$u(z, a) = -za - (a^2 - 1)^2.$$

The principal's problem is

$$\max_{(z, a)} \{v(z, a) : u(z, a) \geq -2 \text{ and } a \in \arg \max_{a'} u(z, a')\}. \quad (1)$$

If we use the FOA, the solutions are $(z, a) = (3/2, 1/2)$ and $(-3/2, -1/2)$, which are not incentive compatible. Thus, the FOA is invalid.

Since this problem is so simple, we can solve it by inspection. We show that $(z, a) = \{(0, 1), (0, -1)\}$ is the set of optimal solutions to (1). Clearly, $a = \pm 1$ is a best response to $z = 0$, providing a utility of -2 for the principal. To show that $z \neq 0$ is not an optimal choice for the principal, observe that for a fixed z , the agent's first-order condition yields

$$a(a^2 - 1) = -z/4, \quad (2)$$

where

$$\operatorname{sgn}(a(a^2 - 1)) = \begin{cases} + & \text{if } a > 1 \text{ or } a \in (-1, 0), \\ - & \text{if } a < -1 \text{ or } a \in (0, 1), \\ 0 & \text{otherwise.} \end{cases}$$

Thus, from (2), if $z > 0$, then the optimal choice of a is either $a < -1$ or $a \in (0, 1)$ (the corner solution $a = 2$ is not optimal, since $\frac{\partial}{\partial a} u(z, 2) < 0$). Also, observe that $a \in (0, 1)$ cannot be optimal, since choosing action $-a$ instead only improves the agent's utility. Hence, an optimal response to $z > 0$ must satisfy $a < -1$. However, this implies that $v(z, a) < -2$, and so $z > 0$ is not an optimal choice of the principal (setting $z = 0$ gives the principal a utility of -2). Nearly identical reasoning shows that $z < 0$ is also not an optimal choice for the principal. This verifies that $(z^*, a^*) = \{(0, 1), (0, -1)\}$ are the optimal solutions to (1). $\diamond$

To handle situations where the FOA is invalid, [Mirrlees \(1999\)](#) recognized that difficulties arise when pairs (w, a) satisfy (FOC(a)) but w fails to implement a . To combat this, Mirrlees reintroduced no-jump constraints from (IC). The resulting problem (cf. [Mirrlees 1986](#)) is

$$\max_{(w, a)} V(w, a) \tag{3a}$$

$$\text{subject to } U(w, a) \geq \underline{U} \tag{3b}$$

$$U_a(w, a) = 0 \tag{3c}$$

$$U(w, a) - U(w, \hat{a}) \geq 0 \quad \text{for all } \hat{a} \text{ such that } U_a(w, \hat{a}) = 0 \tag{3d}$$

(where the complication of corner solutions is ignored for simplicity).⁴ If a candidate contract violates a no-jump constraint in (3d), then an optimizing agent can improve his expected utility by “jumping” to an alternate best response. The *best* choice of alternate action $\hat{a}^*$ given w is included among the no-jump constraints, since such an $\hat{a}^*$ satisfies the first-order condition $U_a(w, \hat{a}^*) = 0$. Hence, if a candidate contract satisfies all no-jump constraints, it must implement a^* . The practical challenge in applying Mirrlees's approach is generating all of the necessary no-jump constraints. In principle, it requires knowing all of the stationary points to the agent's problem for every feasible contract. This enumeration of policies may well be intractable, and no general procedure to systematically produce them is known. However, if additional information can guide the choice of no-jump constraints (such as having a priori knowledge of the optimal contract and its best responses), then Mirrlees's approach can indeed recover the optimal contract. The following example demonstrates this approach and is in the spirit of how Mirrlees illustrated his method.

⁴If corner solutions are considered, (3c) is replaced by (FOC(a)) and instead of (3d), we have one no-jump constraint for every $\hat{a}$ such that (FOC($\hat{a}$)) holds.

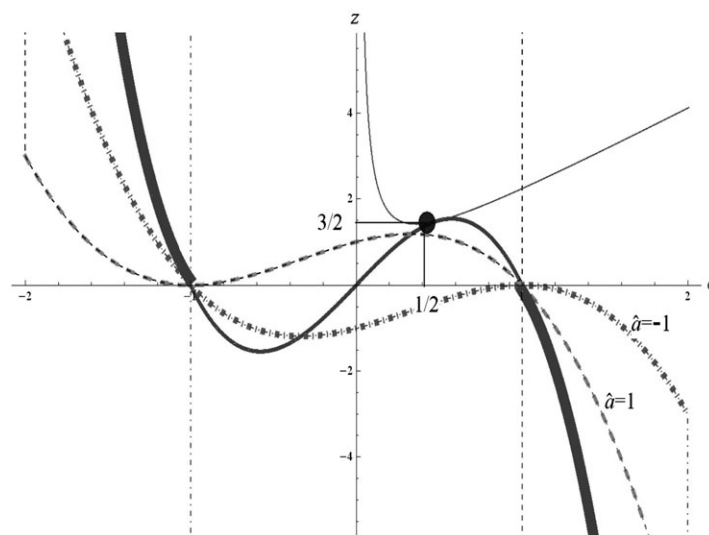


FIGURE 1. Plot for [Example 2](#). The filled curve is the locus of points that satisfy the first-order conditions. The two heavy curve segments form the best response curve. The region below the dotted line for $a \geq -1$ and above the dotted line for $a \leq -1$ captures those points that satisfy $u(z, a) - u(z, -1) \geq 0$. The region above the dashed line for $a \leq 1$ and below the dashed line for $a \geq 1$ captures those points that satisfy $u(z, a) - u(z, 1) \geq 0$. The light curve is the indifference curve of the principal.

EXAMPLE 2 ([Example 1](#) (continued)). If we know a priori the two best responses to an optimal contract, $\hat{a} = 1$ and $\hat{a} = -1$ (as determined in [Example 1](#)), we may solve (1) in the manner

$$\max_{(z,a)} v(z, a)$$

subject to the first-order condition

$$u_a(z, a) = -4a(a^2 - 1) - z = 0$$

and the no-jump constraints

$$u(z, a) - u(z, \hat{a}) \geq 0$$

for $\hat{a} \in \{1, -1\}$. According to (3), we should include many more no-jump constraints, but in fact we show that these two are sufficient to determine the optimal solution. [Figure 1](#) illustrates the constraint sets and optimal solutions.

We plot the first-order condition curve, the best response set, and the regions for the two constraints:

$$\begin{aligned} u(z, a) - u(z, 1) &\geq 0, \\ u(z, a) - u(z, -1) &\geq 0. \end{aligned}$$

are described in the caption of the figure. The region $\{(z, a) : u(z, a) - u(z, \hat{a}) \geq 0\}$ lies below the curve

$$z = -(\hat{a} + a)(\hat{a}^2 + a^2 - 2)$$

for $a > \hat{a}$ and above the curve for $a < \hat{a}$. These constraints preclude the optimal solution of the FOA: $(z, a) = (3/2, 1/2)$ and $(-3/2, -1/2)$. The only contract–action pairs that satisfy these conditions are $(z^*, a^*) = \{(0, 1), (0, -1)\}$, the optimal solutions to (1) (as established in [Example 1](#)). $\diamond$

In our approach, we show how, under additional monotonicity assumptions, reintroducing a single no-jump constraint is all that is required. Moreover, this single constraint can be found by solving an optimization problem in the alternate action $\hat{a}$. The next two sections describe and justify this procedure.

3. THE SANDWICH RELAXATION

We first introduce a family of restrictions on (P) that vary the right-hand side of the (IR) constraint (for reasons that become clear later). Consider the parametric problem

$$\begin{aligned} \max_{w \geq \underline{w}, a \in \mathbb{A}} \quad & V(w, a) \\ \text{subject to} \quad & U(w, a) \geq b \\ & U(w, a) - U(w, \hat{a}) \geq 0 \quad \text{for all } \hat{a} \in \mathbb{A} \end{aligned} \tag{P|b}$$

with parameter $b \geq \underline{U}$. The original problem (P) is precisely (P| $\underline{U}$). We restrict $b \geq \underline{U}$ so that $\text{val}(\text{P|}b) \leq \text{val}(\text{P})$ and a feasible solution of (P| b) remains feasible to (P). We restate (P| b) using an inner minimization over $\hat{a}$. Observe that (P| b) is equivalent to

$$\begin{aligned} \max_{w \geq \underline{w}, a \in \mathbb{A}} \quad & V(w, a) \\ \text{subject to} \quad & U(w, a) \geq b \\ & \inf_{\hat{a} \in \mathbb{A}} \{U(w, a) - U(w, \hat{a})\} \geq 0. \end{aligned} \tag{4}$$

To clarify the relationships between w , a , and $\hat{a}$, we pull the minimization operator out from the constraint (4) and behind the objective function. This requires handling the possibility that a choice of w does not implement the chosen a , in which case (4) is violated. We handle this issue as follows. Given $b \geq \underline{U}$, define the set

$$\mathcal{W}(\hat{a}, b) \equiv \{(w, a) : U(w, a) \geq b \text{ and } U(w, a) - U(w, \hat{a}) \geq 0\}$$

and the characteristic function

$$V^I(w, a | \hat{a}, b) \equiv \begin{cases} V(w, a) & \text{if } (w, a) \in \mathcal{W}(\hat{a}, b), \\ -\infty & \text{otherwise.} \end{cases}$$

This is constructed so that if $V^I(w, a|\hat{a}, b)$ is maximized over (w, a) at a finite objective value, then $(w, a) \in \mathcal{W}(\hat{a}, b)$. Similarly, if maximizing $\inf_{\hat{a} \in \mathbb{A}} V^I(w, a|\hat{a}, b)$ over (w, a) results in a finite objective value, then we know (w, a) lies in $\mathcal{W}(\hat{a}, b)$ for all $\hat{a} \in \mathbb{A}$. This implies that (w, a) is feasible to $(P|b)$, and demonstrates the equivalence of $(P|b)$ and the problem

$$\max_{a \in \mathbb{A}} \max_{w \geq \underline{w}} \inf_{\hat{a} \in \mathbb{A}} V^I(w, a|\hat{a}, b). \quad (\text{Max-Max-Min}|b)$$

We explore what transpires when we swap the order of optimization in $(\text{Max-Max-Min}|b)$ so that $\hat{a}$ is chosen *before* w . That is, we consider

$$\max_{a \in \mathbb{A}} \inf_{\hat{a} \in \mathbb{A}} \max_{w \geq \underline{w}} V^I(w, a|\hat{a}, b),$$

which is equivalent to

$$\max_{a \in \mathbb{A}} \inf_{\hat{a} \in \mathbb{A}} \max_{w \geq \underline{w}} \{V(w, a) : (w, a) \in \mathcal{W}(\hat{a}, b)\}, \quad (\text{SAND}|b)$$

since an optimal choice of a precludes a subsequent optimal choice of $\hat{a}$ that sets $\mathcal{W}(\hat{a}, b) = \emptyset$, implying $V^I(w, a|\hat{a}, b) = V(w, a)$ when w is optimally chosen. We call $(\text{SAND}|b)$ the *sandwich problem* given bound b , where “sandwich” refers to the fact that the minimization over $\hat{a}$ is sandwiched between two maximizations.

Our method allows for the nonexistence of a minimizer to the inner minimization over $\hat{a}$. However, the next lemma shows that the outer maximization over a always possesses a maximizer. This follows by establishing the upper semicontinuity of the value function over the inner two optimization problems.

LEMMA 1. *There exists a maximizer to the outer maximization problem in $(\text{SAND}|b)$.*

Even when the inner minimization over $\hat{a}$ does not exist, we call (a^*, w^*) , where $V(w^*, a^*) = \text{val}(\text{SAND}|b)$ is an optimal solution to $(\text{SAND}|b)$. If the inner minimization is attained at an action $\hat{a}^*$, then we can say that $(a^*, \hat{a}^*, w^*)$ is an optimal solution to $(\text{SAND}|b)$ without confusion.

LEMMA 2. *For every $b \geq \underline{U}$, $\text{val}(P|b) \leq \text{val}(\text{SAND}|b)$. Moreover, if there exists an optimal solution (w^*, a^*) to (P) such that $U(w^*, a^*) \geq b$, then $\text{val}(P) \leq \text{val}(\text{SAND}|b)$.*

From Lemma 2, we are justified in calling $(\text{SAND}|b)$ the *sandwich relaxation* of $(P|b)$. There are two related benefits to studying the sandwich relaxation. First, changing the order of optimization from Max-Max-Min to Max-Min-Max improves analytical tractability. The map that describes which optimal contracts support a given action a against deviation to a specific alternate best response $\hat{a}$ has desirable topological properties. These properties can be used to determine the “minimizing” alternative best response without resorting to enumeration, as is required in the worst-case in Mirrlees’s approach. By contrast, to solve the original problem $(\text{Max-Max-Min}|b)$, one must work with the best-response set $a^{\text{BR}}(w)$ as a constraint for the inner maximization over w . The best-response set is notoriously ill-structured. More details are provided in Section 3.1.

Second, if b satisfies a property called tightness-at-optimality (defined below) and other sufficient conditions are met, the sandwich relaxation is *equivalent* to (P). More details are provided in Section 3.2.

3.1 Analytical benefit of changing the order of optimization

By changing the order of optimization, we solve for an optimal contract w *given* a choice between implementable action a and alternate best response $\hat{a}$. The resulting problem is

$$\max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U(w, a) - U(w, \hat{a}) \geq 0\}. \quad (\text{SAND}|a, \hat{a}, b)$$

We show that this problem has a unique solution, and provide necessary and sufficient optimality conditions.

The approach is to use Lagrangian duality. The Lagrangian function of (SAND| $a, \hat{a}, b$) is

$$\mathcal{L}(w, \lambda, \delta|a, \hat{a}, b) = V(w, a) + \lambda[U(w, a) - b] + \delta[U(w, a) - U(w, \hat{a})], \quad (5)$$

where $\lambda \geq 0$ and $\delta \geq 0$ are the multipliers for $U(w, a) \geq b$ and $U(w, a) - U(w, \hat{a}) \geq 0$, respectively. The Lagrangian dual is

$$\inf_{\lambda, \delta \geq 0} \max_{w \geq \underline{w}} \mathcal{L}(w, \lambda, \delta|a, \hat{a}, b). \quad (6)$$

Consider the inner maximization problem of (6) over w . By Assumption A1.4, we can express the Lagrangian (5) as

$$\mathcal{L}(w, \lambda, \delta|a, \hat{a}, b) = \int L(w(x), \lambda, \delta|x, a, \hat{a}, b) f(x, a) dx,$$

where $L(\cdot, \cdot, \cdot|x, a, \hat{a}, b)$ is a function from $\mathbb{R}^3 \rightarrow \mathbb{R}$ with

$$\begin{aligned} L(y, \lambda, \delta|x, a, \hat{a}, b) &= v(\pi(x) - y) + \lambda(u(y) - c(a) - b) \\ &\quad + \delta \left[u(y) \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) - c(a) + c(\hat{a}) \right] \\ &= v(\pi(x) - y) + \left[\lambda + \delta \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) \right] u(y) \\ &\quad - \lambda(c(a) + b) - \delta(c(a) - c(\hat{a})), \end{aligned}$$

where the ratio $1 - f(x, \hat{a})/f(x, a)$ results from factoring $f(x, a)$ from the terms involving u . This is possible since $f(\cdot, a)$ has the same support for all a .

The inner maximization of $\mathcal{L}(w, \lambda, \delta|a, \hat{a}, b)$ over w in (6) can be done pointwise via

$$\max_{y \geq \underline{w}} L(y, \lambda, \delta|x, a, \hat{a}, b) \quad (7)$$

for each x and setting $w(x) = y$, where y is an optimal solution to (7). Two cases can occur. If $\lambda + \delta(1 - f(x, \hat{a})/f(x, a)) \leq 0$, then $L(y, \lambda, \delta|x, a, \hat{a}, b)$ is a decreasing function of y by Assumptions A1.9 and A1.10. Hence, the unique optimal solution to (7) is $y = \underline{w}$.

Alternatively, if $\lambda + \delta(1 - f(x, \hat{a})/f(x, a)) > 0$, then $L(y, \lambda, \delta|x, \hat{a})$ is strictly concave in y (again by Assumptions A1.9 and A1.10). If $\frac{\partial}{\partial y}L(\underline{w}, \lambda, \delta|x, a, \hat{a}, b) \leq 0$, then the corner solution $y = \underline{w}$ is optimal; otherwise there exists a unique y such that $\frac{\partial}{\partial y}L(y, \lambda, \delta|x, a, \hat{a}, b) = 0$ holds. In both cases, (7) has a unique optimal solution $w(x)$. Hence, the optimal solution $w: \mathcal{X} \rightarrow \mathbb{R}$ to the inner maximization of (6) satisfies

$$w(x) \begin{cases} \text{solves } \frac{\partial}{\partial y}L(w(x), \lambda, \delta|x, a, \hat{a}, b) = 0 \\ \text{if } \lambda + \delta\left(1 - \frac{f(x, \hat{a})}{f(x, a)}\right) > 0 \text{ and } \frac{\partial}{\partial y}L(\underline{w}, \lambda, \delta|x, a, \hat{a}, b) > 0, \\ = \underline{w} \text{ otherwise.} \end{cases}$$

Expressing the derivatives and dividing by $u'(w(x))$ (which is valid since $u' > 0$ by Assumption A1.9) yields

$$w(x) \begin{cases} \text{solves } \frac{v'(\pi(x) - w(x))}{u'(w(x))} = \lambda + \delta\left(1 - \frac{f(x, \hat{a})}{f(x, a)}\right) \\ \text{if } \frac{v'(\pi(x) - \underline{w})}{u'(\underline{w})} < \lambda + \delta\left(1 - \frac{f(x, \hat{a})}{f(x, a)}\right), \\ = \underline{w} \text{ otherwise.} \end{cases} \quad (8)$$

Since v' and u' are both positive, the condition $v'(\pi(x) - \underline{w})/u'(\underline{w}) < \lambda + \delta(1 - f(x, \hat{a})/f(x, a))$ implies $\lambda + \delta(1 - f(x, \hat{a})/f(x, a)) > 0$, handling both cases detailed above.

We just showed that, given $(\lambda, \delta, a, \hat{a}, b)$, there is a unique choice w , denoted $w_{\lambda, \delta}(a, \hat{a}, b)$, that satisfies (8). Such contracts are significant for our analysis and warrant a formal definition.

DEFINITION 1. Any contract that satisfies (8) for some choice of $(\lambda, \delta, a, \hat{a}, b)$ is called a *generalized Mirrlees–Holmstrom* (GMH) contract. These contracts are generalized versions of Mirrlees–Holmstrom contracts known for the special case of a binary action.

Observe that GMH contracts are continuous in x . There are five parameters $(\lambda, \delta, a, \hat{a}, b)$ in a GMH contract. However, Lemma 3 below shows that each GMH contract is a function of only three parameters: $a, \hat{a}$, and b .

LEMMA 3. Suppose Assumptions 1–3 hold. For every $(a, \hat{a}, b)$ with $\hat{a} \neq a$, there exist unique Lagrangian multipliers λ^* and δ^* and a unique contract w^* such that the following statements hold:

- (i) The variable w^* satisfies (8) for λ^* and δ^* (in particular, w^* is a GMH contract).
- (ii) Strong duality between (SAND| $a, \hat{a}, b$) and (6) holds and, in particular, the complementary slackness conditions

$$\lambda^* \geq 0, \quad U(w^*, a) - b \geq 0, \quad \text{and} \quad \lambda^*[U(w^*, a) - b] = 0, \quad (\text{ii-a})$$

$$\delta^* \geq 0, \quad U(w^*, a) - U(w^*, \hat{a}) \geq 0, \quad \text{and} \quad \delta^*[U(w^*, a) - U(w^*, \hat{a})] = 0 \quad (\text{ii-b})$$

are satisfied.

Moreover, the following additional properties hold:

- (iii) The equality $(\lambda^*, \delta^*) = (\lambda(a, \hat{a}, b), \delta(a, \hat{a}, b))$ is an upper semicontinuous function of $(a, \hat{a}, b)$, and is continuous and differentiable at any $(a, \hat{a}, b)$ where $a \neq \hat{a}$.
- (iv) The equality $w^* = w_{\lambda(a, \hat{a}, b), \delta(a, \hat{a}, b)}(a, \hat{a}, b)$ is an upper semicontinuous function of $(a, \hat{a}, b)$, and is continuous and differentiable at any $(a, \hat{a}, b)$ where $a \neq \hat{a}$.

Lemma 3(iv) leaves open the possibility that there is a jump discontinuity when $a = \hat{a}$. As an illustration, consider the case where the principal is risk-neutral and the FOA is valid. When $\hat{a} > a$, the optimal solution to (SAND| $a, \hat{a}, b$) is the first-best contract. However, as $\hat{a} - a \rightarrow 0^-$, we have

$$\lim_{\hat{a} - a \rightarrow 0^-} V(w_{\lambda(a, \hat{a}, b), \delta(a, \hat{a}, b)}(a, \hat{a}, b), a) = \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U_a(w, a) = 0\} \\ < \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b\}.$$

Therefore, the value function is not continuous at that point.⁵

Lemma 3 provides insight into the inner inf of (SAND| b). Given an $a \in \mathbb{A}$, suppose the infimizing sequence $\hat{a}^n$ to the inner inf converges to some a' . If $a' \neq a$, then, in fact, the infimum is attained by the continuity of w^* from Lemma 3(iv). An issue arises if $a' = a$ and the infimum is not attained, since this is a point of discontinuity of w^* . The following result analyzes this scenario. We also refer the reader to Section 5 below, which provides additional discussion of this case.

LEMMA 4. *If the minimization of $\inf_{\hat{a}} \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U(w, a) - U(w, \hat{a}) \geq 0\}$ is not attained, then*

$$\inf_{\hat{a}} \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U(w, a) - U(w, \hat{a}) \geq 0\} \\ = \max_w \{V(w, a) : U(w, a) \geq b, (FOC(a))\}, \quad (9)$$

where $(FOC(a))$ is as defined in Section 2.2.

This result shows that when the infimum is not attained for a given action a , it suffices to take a “first-order approach” locally at a .

3.2 Tightness of the sandwich relaxation

The previous subsection provides a toolbox for analyzing the sandwich relaxation (SAND| b). However, there remains the question of whether that relaxation is worth solving at all. In particular, we may ask whether there exists a b that makes (SAND| b) a *tight* relaxation; i.e., whether an optimal solution (a^*, w^*) to (SAND| b) yields an optimal solution (w^*, a^*) to (P). The following example illustrates a situation where such a choice is possible.

⁵We thank an anonymous referee for alerting us to this observation.

EXAMPLE 3 (Example 1 (continued)).

We solve the sandwich relaxation (SAND|0) and show that (SAND|0) is a tight relaxation.⁶ That is, we solve

$$\max_{a \in [-2, 2]} \inf_{\hat{a} \in [-2, 2]} \max_z \{v(z, a) : u(z, a) \geq 0 \text{ and } u(z, a) - u(z, \hat{a}) \geq 0\}, \quad (10)$$

where

$$v(z, a) = za - 2a^2 \quad \text{and} \quad u(z, a) = -za - (a^2 - 1)^2.$$

We break up the outermost optimization (over a) across two subregions of $[-2, 0]$ and $[0, 2]$. The optimal value of (10) can be found by taking the larger of the two values across the two subregions. We consider $a \in [0, 2]$ first. In this case $v(z, a)$ is increasing in z and thus $\hat{a}$ is chosen to minimize z . We show how z relates to the choice of a and $\hat{a}$. The $u(z, a) \geq 0$ constraint cannot be satisfied when $a = 0$ and so it is equivalent to

$$z \leq -\frac{(a^2 - 1)^2}{a}, \quad (11)$$

since dividing by $a \neq 0$ is legitimate. The no-jump constraint $u(z, a) - u(z, \hat{a}) \geq 0$ is equivalent to

$$z \begin{cases} \geq -(\hat{a} + a)(\hat{a}^2 + a^2 - 2) & \text{for } \hat{a} > a, \\ \leq -(\hat{a} + a)(\hat{a}^2 + a^2 - 2) & \text{for } \hat{a} < a, \\ \in (-\infty, \infty) & \text{for } \hat{a} = a. \end{cases} \quad (12)$$

Clearly, $\hat{a} = a$ will never be chosen in the inner minimization over $\hat{a}$ in (10) since it cannot prevent sending $z \rightarrow \infty$, when the goal is to minimize z . When $\hat{a} > a$, observe that

$$\begin{aligned} & \inf_{\hat{a} > a} -(\hat{a} + a)(\hat{a}^2 + a^2 - 2) \\ &= \begin{cases} 4a - 4a^3 & \text{for } 1/\sqrt{3} \leq a \leq 2, \\ \frac{4}{27}(9a - 5a^3) + \frac{4}{27}\sqrt{2}\sqrt{(3 - a^2)^3} & \text{for } 0 \leq a \leq 1/\sqrt{3}. \end{cases} \end{aligned} \quad (13)$$

When $a \in [0, 1)$, one can verify that

$$\inf_{\hat{a} > a} -(\hat{a} + a)(\hat{a}^2 + a^2 - 2) > 0 > -(a^2 - 1)^2/a$$

using (13). By (12), this implies $z > (a^2 - 1)^2/a$ when $\hat{a} > a$, violating (11). Hence, when $a \in [0, 1)$, the inner minimization over $\hat{a}$ in (10) will choose $\hat{a} > a$ and thus make a choice of z infeasible. This drives the value of the inner minimization over $\hat{a}$ to $-\infty$. This, in turn, implies that $a \in [0, 1)$ will never be chosen in the outer maximization, and so we may restrict attention to $a \in [1, 2]$.

⁶In fact, one can show that setting $b = \underline{U} = -2$ does not give rise to a tight relaxation. For details, see the discussion following (30) below.

When $a \in [1, 2]$, we return to (12) and consider the two cases (i) $\hat{a} > a$ and (ii) $\hat{a} < a$. In case (i), note that

$$\inf_{\hat{a} > a} -(\hat{a} + a)(\hat{a}^2 + a^2 - 2) = 4a - 4a^3 \leq -\frac{(a^2 - 1)^2}{a},$$

when $a \in [1, 2]$ and so from (11)–(13) we have

$$4a - 4a^3 \leq z \leq -\frac{(a^2 - 1)^2}{a}.$$

However, in case (ii), we have from (11) and (12) that

$$z \leq \min \left\{ \frac{(a^2 - 1)^2}{a}, \inf_{\hat{a} < a} -(\hat{a} + a)(\hat{a}^2 + a^2 - 2) \right\}.$$

Note that

$$\inf_{\hat{a} < a} -(\hat{a} + a)(\hat{a}^2 + a^2 - 2) = 4a - 4a^3 \quad \text{for } 1 \leq a \leq 2 \quad (14)$$

and $4a - 4a^3 < -(a^2 - 1)^2/a$ when $a \in [1, 2]$. Observe that the infimum is not attained since the only real solution to $-(\hat{a} + a)(\hat{a}^2 + a^2 - 2) = 4a - 4a^3$ when $a \in [1, 2]$ is $\hat{a} = a$. Lemma 4 applies and yields

$$z^*(a) = 4a - 4a^3 \quad (15)$$

via (14). Since the principal's utility $v(z^*(a), a)$ is decreasing in $a \in [1, 2]$, we obtain the solution $a^* = 1$ and the optimal choice of z^* is thus $z^*(1) = 0$. One can see this graphically in Figure 2.⁷

We return to the case where $a \in [-2, 0]$. Nearly identical reasoning (with care to adjust negative signs) shows $a^* = -1$ and, again, the optimal choice of z is $z^*(1) = 0$. Hence, the overall problem (10) gives rise to *two* optimal choices of (z^*, a^*) , namely $(0, 1)$ and $(0, -1)$. However, this is precisely the optimal solution to the original problem (1), as shown by inspection in Example 1. $\diamond$

Note that by choosing b correctly in the above example we were able to arrive at the first-order condition curve $U_a(z, a) = 0$ used in Mirrlees's approach. This underscores that we do not need to *explicitly* include the FOC in our definition of the sandwich relaxation. This issue is taken up more carefully in Section 5. Comparing Figure 1 and Figure 2, we see that the (IR) is not needed to specify the optimal contract in Figure 1, but is needed (with an adjusted right-hand side) when using the sandwich relaxation in Figure 2. However, the first-order condition curve does not appear in Figure 2 to characterize the optimal contract.

⁷This example has the special structure that the FOA applies *locally*. That is, given an a , the optimal choice of z is uniquely determined by the first-order condition. The classic example in Mirrlees (1999) also has this property.

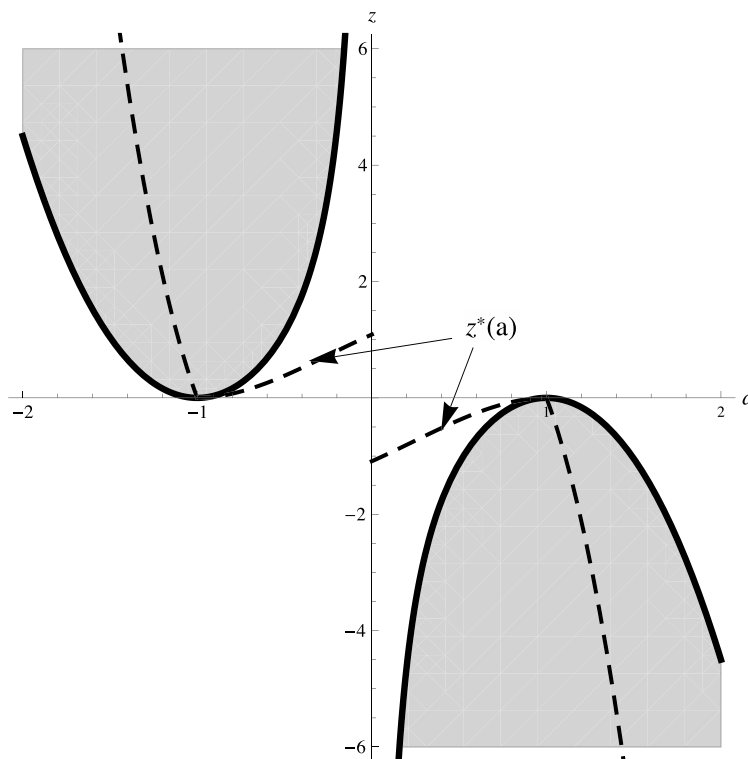


FIGURE 2. Plot for Example 3. The non-dashed curve and filled region are those (z, a) that satisfy the constraint $U(z, a) \geq 0$. The dashed curves are those (z, a) that satisfy the inner maximization over z given by (15). Observe that the optimal solution in the region $a \in [0, 2]$ is $(z, a) = (0, 1)$ since the principal's utility is increasing in z .

Of course, the question remains as to whether there always exists a b such that $(\text{SAND}|b)$ is a tight relation of (P) , and if so, how to determine it. We state the following definition.

DEFINITION 2. We say $b \geq \underline{U}$ is *tight at optimality* (or simply *tight*) if there exists an optimal solution (w^*, a^*) to (P) such that $b = U(w^*, a^*)$.

At least one such b exists by Assumption 3. The main result of this section is to show that for such a b , the sandwich relaxation $(\text{SAND}|b)$ is tight under certain conditions. The key assumption is that f satisfies the *monotone likelihood ratio property* (MLRP), where, for any a , $\frac{\partial \log f(\cdot, a)}{\partial a}$ is nondecreasing. This property is well known in the literature (see Hölmstrom 1979, Rogerson 1985, and others).

ASSUMPTION 4. *The output distribution f satisfies the MLRP condition.*

The following theorem is the key technical result of the paper.

THEOREM 1. *Suppose Assumptions 1–4 hold. If b is tight at optimality, then $(\text{SAND}|b)$ is a tight relaxation; that is, $\text{val}(\text{SAND}|b) = \text{val}(\text{P})$ and if $(a^\#, \hat{a}^\#, w^\#)$ is an optimal solution to $(\text{SAND}|b)$, then $(w^\#, a^\#)$ is an optimal solution to (P) . If the infimum in $(\text{SAND}|b)$ is not attained, and $(a^\#, w^\#)$ is an optimal solution to the inner and outer maximization in $(\text{SAND}|b)$, then $(w^\#, a^\#)$ is an optimal solution to (P) .*

The proof of [Theorem 1](#) is involved and relies on several nontrivial, but largely technical, intermediate results. Full details are found in the [Appendix](#), along with further discussion. We note that [Lemma 4](#) is essential for the case where the infimum is not attained.

For the sake of developing intuition regarding the proof of [Theorem 1](#), we consider here the special case where $\mathcal{X}$ is a singleton and the inner infimum is attained. Of course, the single-outcome case is not a difficult problem to solve and provides little economic intuition, but it does highlight some important features of the more general argument that we discuss below.

When $\mathcal{X}$ is a singleton, contracts w are characterized by a single number $z = w(x_0)$ (following the notation of [Example 2](#) and [Mirrlees 1999](#)), and so $U(w, a) = u(z) - c(a)$ and $V(w, a) = v(\pi(x_0) - z)$. For consistency, we denote the minimum wage by $\underline{z}$ (as opposed to $\underline{w}$).

PROOF OF THEOREM 1 FOR A SINGLE-DIMENSIONAL CONTRACT. Since b is tight at optimality, there exists an optimal solution (z^*, a^*) of (P) such that $b = U(z^*, a^*)$. Let $(a^\#, \hat{a}^\#, z^\#)$ be an optimal solution to $(\text{SAND}|b)$.

There are two cases to consider.

Case 1: $U(z^\#, a^\#) = b$. By [Lemma 2](#), we know that $\text{val}(\text{P}) \leq \text{val}(\text{SAND}|b)$. It suffices to argue that $\text{val}(\text{SAND}|b) \leq \text{val}(\text{P})$. By the optimality of $(a^\#, \hat{a}^\#, z^\#)$ in $(\text{SAND}|b)$, we know that

$$V(z^\#, a^\#) = \inf_{\hat{a} \in \mathbb{A}} \max_{z \geq \underline{z}} \{V(z, a^\#) : U(z, a^\#) \geq b, U(z, a^\#) - U(z, \hat{a}) \geq 0\}. \quad (16)$$

Let $\hat{a}'$ be a best response to $z^\#$. Then from the minimization over $\hat{a}$ in (16), we have

$$V(z^\#, a^\#) \leq \max_{z \geq \underline{z}} \{V(z, a^\#) : U(z, a^\#) \geq b, U(z, a^\#) - U(z, \hat{a}') \geq 0\}. \quad (17)$$

Suppose (17) holds with equality. Since V is decreasing in z (under [Assumption A1.10](#)) and the feasible region is single-dimensional, the optimal solution to the right-hand side problem is unique and, therefore, $z^\#$ must be that unique optimal solution under the equality assumption. This implies $z^\#$ is feasible to the right-hand side problem and so $U(z^\#, a^\#) \geq U(z^\#, \hat{a}')$. Since $\hat{a}'$ is a best response to $z^\#$, $a^\#$ is also. This implies that $(z^\#, a^\#)$ is a feasible solution to (P) . Thus, $\text{val}(\text{SAND}|b) \leq \text{val}(\text{P})$, establishing the result.

Hence, it remains to argue that (17) is satisfied with equality. Suppose otherwise that

$$V(z^\#, a^\#) < \max_{z \geq \underline{z}} \{V(z, a^\#) : U(z, a^\#) \geq b, U(z, a^\#) - U(z, \hat{a}') \geq 0\}.$$

There must exist a z' in the argmax of the right-hand side such that $V(z^\#, a^\#) < V(z', a^\#)$. Since V is strictly decreasing in z , this implies $z^\# > z'$. However, since U is increasing in z , this further implies that $U(z', a^\#) < U(z^\#, a^\#) = b$ (where the equality holds under the assumption of Case 1). That is, $U(z', a^\#) < b$, contradicting the feasibility of z' to $(\text{SAND}|b)$.

Case 2: $U(z^\#, a^\#) > b$. This case requires the following intermediate lemma, whose proof is provided in the [Appendix](#).

LEMMA 5. *Let $(a^\#, z^\#)$ be an optimal solution to the single-dimensional version of $(\text{SAND}|b)$ with $U(z^\#, a^\#) > b$ (in particular, the infimum in $(\text{SAND}|b)$ need not be attained). Then there exists an $\epsilon > 0$ such that the perturbed problem $(\text{SAND}|b + \epsilon)$ also has an optimal solution $(a_\epsilon^\#, z_\epsilon^\#)$ with $U(z_\epsilon^\#, a_\epsilon^\#) = b + \epsilon$ and the same optimal value; that is, $V(z_\epsilon^\#, a_\epsilon^\#) = V(z^\#, a^\#) = \text{val}(\text{SAND}|b)$.*

The proof of this lemma relies on strong duality and the fact that if a constraint is slack, the dual multiplier on that constraint is 0 by complementary slackness. A small perturbation of the right-hand side of a slack constraint does not impact the optimal value. This argument is standard (see, for instance, [Bertsekas 1999](#)) in the absence of the inner minimization problem $\inf_{\hat{a}}$ in $(\text{SAND}|b)$. With the inner minimization, the proof becomes nontrivial.

Returning to our proof of Case 2, by [Lemma 5](#) there exists an $\epsilon > 0$ and an optimal solution $(a_\epsilon^\#, z_\epsilon^\#)$ to $(\text{SAND}|b + \epsilon)$, where $U(z_\epsilon^\#, a_\epsilon^\#) = b + \epsilon$ and $\text{val}(\text{SAND}|b + \epsilon) = \text{val}(\text{SAND}|b)$. Apply the logic of Case 1 to the problem $(\text{SAND}|b + \epsilon)$ and conclude that $\text{val}(\text{SAND}|b + \epsilon) = \text{val}(\text{P})$. Hence, since $\text{val}(\text{SAND}|b + \epsilon) = \text{val}(\text{SAND}|b)$, $(\text{SAND}|b)$ is a tight relaxation of (P) . $\square$

We provide here some intuition behind [Theorem 1](#) in the single-outcome setting. For a given target action a^* , we can think of the contracting problem as a sequential game where the principal chooses z and the agent chooses $\hat{a}$. The original (IC) constraint is equivalent to the situation that the principal chooses z first, followed by the agent's choice of $\hat{a}$. So the optimal choice of z should take all possible $\hat{a}$ into consideration. The agent has a second-mover advantage. Now consider a change in the order of decisions and let the agent choose $\hat{a}$ first, with the principal choosing z in response. In this case, the principal has a second-mover advantage since the principal need not consider all possible $\hat{a}$. This provides intuition behind the bound in [Lemma 2](#). However, if the agent utility bound b is tight given a^* , the principal cannot gain an advantage by moving second. No choice of contract by the principal can drive the agent's utility down any further. Since the principal and agent have a direct conflict of interest over the direction of z , this means the principal cannot improve her utility. In other words, the order of decisions does not matter when b is tight and so the sandwich problem provides a tight relaxation.

This argument relies on the fact that w is unidimensional. In the multidimensional case, we parameterize the payment function through a unidimensional z using a variational argument. As long as a conflict of interest exists, we obtain a similar intuition and result. An analogous result to [Lemma 5](#) is also leveraged in the argument.

We remark that [Assumption 4](#) is not used in the proof of [Theorem 1](#) for the singleton case. However, [Assumption 4](#) is essential for continuous outcome sets. The MLRP is essential for showing that optimal solutions to sandwich relaxations are, in fact, GMH contracts as defined in [Section 3.1](#). In particular, monotonicity of the output function greatly simplifies the first-order conditions of (P) to reduce them to the necessary and sufficient conditions of (8).

Of course, there remains the question of how to find a tight b . The simplest case is when the reservation utility $\underline{U}$ itself is tight. The following proposition gives a sufficient condition for this to hold.

PROPOSITION 1. *Suppose Assumptions 1–3 hold. Then the reservation utility $\underline{U}$ is tight at optimality if there exist an optimal solution w^* to (P) and an $\delta > 0$ such that $w^*(x) > \underline{w} + \delta$ for almost all $x \in \mathcal{X}$.*

A main task of the next section is to provide a systematic approach to finding a b that is tight at optimality even when the conditions of [Proposition 1](#) fail to hold. We should remark that it is not uncommon in the FOA literature to focus on the case where the limited liability constraint is not binding (see [Rogerson 1985](#) and [Jewitt et al. 2008](#)). If that convention is taken here, [Proposition 1](#) is useful in determining optimal contracts.

4. THE SANDWICH PROCEDURE

The remaining steps to systematically solve (P) are (i) finding a b that is tight at optimality and (ii) determining a systematic way to solve (SAND| b). We approach both tasks concurrently using what we call the *sandwich procedure*. The basic logic of the procedure is to use backward induction, leveraging [Lemma 3](#) and the GMH structure (see [Definition 1](#)) of optimal solutions to (SAND| $a, \hat{a}, b$). The structure of these optimal solutions is used to compute a tight b by solving a carefully designed optimization problem.

THE SANDWICH PROCEDURE

Step 1: Characterize contract. Characterize an optimal solution to the inner maximization in (SAND| b),

$$\max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U(w, a) - U(w, \hat{a}) \geq 0\}, \quad (\text{SAND}|a, \hat{a}, b)$$

as a function of $a \in \mathbb{A}$, $\hat{a} \in \mathbb{A}$ and $b \geq \underline{U}$, where $\hat{a} \neq a$. Denote the resulting optimal contract by $w(a, \hat{a}, b)$.

Step 2: Characterize actions. Determine optimal solutions to the outer maximization and minimization,

$$\max_{a \in \mathbb{A}} \inf_{\hat{a} \in \mathbb{A}} V(w(a, \hat{a}, b), a), \quad (18)$$

as functions of b . If a minimizer $\hat{a}(a, b)$ exists, find $a(b) \in \arg\max_{a \in \mathbb{A}} V(w(a, \hat{a}(a, b), b), a)$ and set $w(b) = w(a(b), \hat{a}(a(b), b), b)$. If a minimizer does not exist,

solve

$$\max_{a \in \mathbb{A}} \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, (\text{FOC}(a))\},$$

which uses (9) from Lemma 4. Call the resulting solution $(a(b), w(b))$.

Step 3: Compute a tight bound. Solve the one-dimensional optimization problem:

$$b^* \equiv \min \left\{ \underset{b \geq \underline{U}}{\text{argmin}} \left\{ V(w(b), a(b)) - \max_{a \in a^{\text{BR}}(w(b))} V(w(b), a) \right\} \right\}. \quad (19)$$

Let $a^* \equiv a(b^*)$, $\hat{a}^* \equiv \hat{a}(a^*, b^*)$ (when it exists), and $w^* \equiv w(b^*)$.

PROPOSITION 2. *For a given b , let $a(b)$, $\hat{a}(a(b), b)$ (if it exists), and $w(b)$ be as defined at the end of Step 2 of the sandwich procedure. Then $(a(b), \hat{a}(a(b), b), w(b))$ is an optimal solution to the sandwich relaxation (SAND| b). If $\hat{a}(a(b), b)$ does not exist, then $(a(b), w(b))$ (as defined in Step 2) solves (SAND| b).*

The proof is essentially by definition and thus is omitted. However, to *guarantee* the tractability of each step, we must make Assumptions 1–4. These same conditions ensure that (SAND| b) is, in fact, a tight relaxation.

THEOREM 2. *Suppose Assumptions 1–4 hold, and let b^* , a^* , and w^* be as defined in Step 3 of the sandwich procedure. Then b^* is tight at optimality, (w^*, a^*) is an optimal solution to (P), and $\text{val}(\text{SAND}|b^*) = \text{val}(P)$.*

Note that if a given b is known to be tight at optimality through some independent means, Step 3 of the procedure can be avoided. A special case of this is when the reservation utility $\underline{U}$ itself is tight at optimality. Proposition 1 gives a sufficient conditions for this to hold. In practice, one may try to solve (SAND| $\underline{U}$) and if the resulting optimal solution $(a^\#, w^\#)$ is implementable (i.e., $w^\#$ implements $a^\#$), then the complete sandwich procedure can be avoided.

In the remainder of this subsection, we provide lemmas that justify each step of the sandwich procedure. This culminates in a proof of Theorem 2 that is relatively straightforward given the work up to that point. In the final subsection, we note that even when all of Assumptions 1–4 do not hold, we can sometimes use the sandwich procedure to construct an optimal contract. We use our motivating example to illustrate how this can be done.

4.1 Analysis of Step 1

We undertake an analysis of this step under Assumptions 1–3 following from Lemma 3 in Section 3.1. The optimal contract $w(a, \hat{a}, b)$ sought in Step 1 is precisely the unique optimal contract guaranteed by Lemma 3(i). That lemma also guarantees that $w(a, \hat{a}, b)$ is a well behaved function of $(a, \hat{a}, b)$.

Indeed, by strong duality (Lemma 3(ii)), the optimal value of (SAND| $a, \hat{a}, b$) is

$$\text{val}(\text{SAND}|a, \hat{a}, b) = \inf_{\lambda, \delta \geq 0} \max_{w \geq \underline{w}} \mathcal{L}(w, \lambda, \delta|a, \hat{a}, b) = \mathcal{L}^*(a, \hat{a}|b),$$

where

$$\mathcal{L}^*(a, \hat{a}|b) \equiv \mathcal{L}(w(a, \hat{a}, b), \lambda(a, \hat{a}, b), \delta(a, \hat{a}, b)|a, \hat{a}, b) \quad (20)$$

is called the *optimized Lagrangian* for the sandwich relaxation. The following result, a straightforward consequence of the theorem of maximum and Lemma 3, shows that the optimized Lagrangian has useful structure for Step 2 of the procedure.

LEMMA 6. *The optimized Lagrangian $\mathcal{L}^*(a, \hat{a}|b)$ is upper semicontinuous. Moreover, it is continuous and differentiable when $a \neq \hat{a}$.*

4.2 Analysis of Step 2

The case where the inner infimum is not attained is sufficiently handled by Lemma 4 and existing knowledge of the FOA. Here we examine the case where the inner infimum is attained, and we provide necessary optimality conditions for a and $\hat{a}$ to optimize (SAND| b) given the contract $w(a, \hat{a}, b)$ and its associated dual multipliers $\lambda(a, \hat{a}, b)$ and $\delta(a, \hat{a}, b)$. In particular, we solve (18) in Step 2 by solving

$$\max_{a \in \mathbb{A}} \inf_{\hat{a} \in \mathbb{A}} \mathcal{L}^*(a, \hat{a}|b) \quad (21)$$

using the definition of the optimized Lagrangian $\mathcal{L}^*$ in (20). The optimal solution to the outer optimization exists since $\mathbb{A}$ is compact and $\mathcal{L}^*$ is upper semicontinuous (via Lemma 6). Moreover, by the differentiability properties of $\mathcal{L}$ (when $\hat{a} \neq a$), we can obtain the following optimality conditions for solutions of (21).

LEMMA 7. *Suppose a^* and $\hat{a}^*$ solve (21) for a given $b \geq \underline{w}$ with $\hat{a}^* \neq a^*$. The following statements hold:*

(i) *For an interior solution $\hat{a}^* \in (\underline{a}, \bar{a})$,*

$$\frac{\partial}{\partial \hat{a}} \mathcal{L}^*(a^*, \hat{a}^*|b) = -\delta^*(a^*, \hat{a}^*, b) U_a(w(a^*, \hat{a}^*, b), \hat{a}^*) = 0$$

and $U_a(w(a^, \hat{a}^*, b), \hat{a}^*) \geq 0$ (≤ 0) for $\hat{a}^* = \bar{a}$ ($\hat{a}^* = \underline{a}$).*

(ii) *For an interior solution $a^* \in (\underline{a}, \bar{a})$, the right derivative is*

$$\frac{\partial}{\partial a^+} \min_{\hat{a} \in \mathbb{A}} \mathcal{L}^*(a^*, \hat{a}^*|b) \leq 0$$

and left derivative is

$$\frac{\partial}{\partial a^-} \min_{\hat{a} \in \mathbb{A}} \mathcal{L}^*(a^*, \hat{a}^*|b) \geq 0,$$

and $\frac{\partial}{\partial \hat{a}} \mathcal{L}^(a^*, \hat{a}^*|b) \leq 0$ (≥ 0) for $a^* = \underline{a}$ ($a^* = \bar{a}$).*

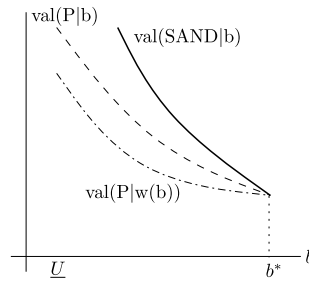


FIGURE 3. An illustration of Step 3 of the sandwich procedure.

Note that the conditions for a^* and $\hat{a}^*$ are not symmetric in (i) and (ii) above. This is because a^* is a function of $\hat{a}^*$ and so has weaker topological properties to leverage for first-order conditions.

4.3 Analysis of Step 3

To work with (19) we reexpress it in a slightly different way. Note that $V(w(b), a(b)) = \text{val}(\text{SAND}|b)$ via Proposition 2. We also denote the optimization problem in the second term inside the argmin of (19) as $(P|w(b))$:

$$\max_{a \in a^{\text{BR}}(w(b))} V(w(b), a). \quad (P|w(b))$$

Thus, we can reexpress (19) as

$$b^* \equiv \min_{b \geq \underline{U}} \{\text{argmin}\{\text{val}(\text{SAND}|b) - \text{val}(P|w(b))\}\}.$$

Note that $(P|w(b))$ is a restriction of $(P|b)$ and so $\text{val}(P|w(b)) \leq \text{val}(P|b) \leq \text{val}(\text{SAND}|b)$ and all three values are decreasing in b . Also from Assumption 3, there exists an optimal solution (w^*, a^*) to (P) and so there exists a b (namely, $b^* = U(w^*, a^*)$) such that all three problems share the same optimal value. Hence, we must have $\min_{b \geq \underline{U}} (\text{val}(\text{SAND}|b) - \text{val}(P|w(b))) = 0$ and so b^* is the first time where $\text{val}(\text{SAND}|b) = \text{val}(P|w(b))$, forcing $\text{val}(\text{SAND}|b) = \text{val}(P|b)$ and implying b^* is tight at optimality. See Figure 3. We make this argument formally in the proof of the following lemma, which also shows that b^* is well defined in the sense that the set $\text{argmin}_{b \geq \underline{U}} \{\text{val}(\text{SAND}|b) - \text{val}(P|w(b))\}$ has a minimum.

LEMMA 8. *If Assumptions 1–4 hold, then there exists a real number b^* that satisfies (19). Furthermore, b^* is tight at optimality.*

4.4 Overall verification of the procedure

We are now ready to prove Theorem 2. The proof is a straightforward application of the lemmas of this section.

PROOF OF THEOREM 2. By Lemma 8, there exists a b^* that satisfies (19) and is tight at optimality. Hence, by Theorem 1, $\text{val}(\text{SAND}|b^*) = \text{val}(P)$ and every optimal solution $(w(b^*), a(b^*))$ to $(\text{SAND}|b^*)$ is optimal to (P) . Note that we need not require that

the infimum is attained. However, when $\hat{a}^*$ is attained with $\hat{a}^* \neq a^*$, the GMH contract $w(a(b^*), \hat{a}(b^*), b^*)$ resulting from [Lemma 3](#) is precisely one such optimal contract where $a(b^*)$ and $\hat{a}(b^*)$ satisfy the optimality conditions of [Lemma 7](#). $\square$

4.5 An illustrative example

Our motivating example serves to illustrate the steps of the sandwich procedure and how to work with (19), even when [Theorem 2](#) does not apply.

EXAMPLE 4 ([Example 1](#) (continued)). Recall that our problem is to solve

$$\max_{(z,a)} \left\{ v(z, a) : u(z, a) \geq -2 \text{ and } a \in \arg \max_{a'} u(z, a') \right\},$$

where $v(z, a) = za - 2a^2$ and $u(z, a) = -za - (a^2 - 1)^2$. We apply each step of the procedure and determine an optimal contract. There is some overlap of analysis from [Example 3](#), but our approach here is more systematic and follows the reasoning of the sandwich procedure.

Step 1: Characterize contract. First, we characterize the optimal solutions $z(a, \hat{a}, b)$ of

$$\max_z \left\{ v(z, a) : u(z, a) \geq b, u(z, a) - u(z, \hat{a}) \geq 0 \right\}, \quad (22)$$

where $a \in [0, 2]$. The case where $a \in [-2, 0]$ is symmetric and analogous reasoning holds throughout. Observe that $v(z, a)$ is increasing in z for fixed a and $\hat{a}$, and so (22) can be solved by simply maximizing z . The constraints on z are (from $u(z, a) \geq b$)

$$z \leq Q(a, b) \quad (23)$$

when $a \neq 0$, where $Q(a, b) \equiv -(b + (a^2 - 1)^2)/a$, and (from $u(z, a) - u(z, \hat{a}) \geq 0$)

$$z \begin{cases} \geq R(a, \hat{a}) & \text{if } \hat{a} > a, \\ \leq R(a, \hat{a}) & \text{if } \hat{a} < a, \\ \in (-\infty, \infty) & \text{if } \hat{a} = a, \end{cases} \quad (24)$$

where $R(a, \hat{a}) \equiv -(\hat{a} + a)(\hat{a}^2 + a^2 - 2)$. Maximizing z subject to (23) and (24) yields

$$\begin{aligned} & z(a, \hat{a}, b) \\ = & \begin{cases} \min\{Q(a, b), R(a, \hat{a})\} & \text{if } (a \neq 0) \wedge (\hat{a} < a), \\ Q(a, b) & \text{if } (a \neq 0) \wedge ((\hat{a} = a) \vee ((\hat{a} > a) \wedge [Q(a, b) \geq R(a, \hat{a})])), \\ R(a, \hat{a}) & \text{if } (a = 0) \wedge (b \leq -1) \wedge (\hat{a} < a), \\ +\infty & \text{if } (a = 0) \wedge (b \leq -1) \wedge (\hat{a} \geq a), \\ -\infty & \text{if } (a \neq 0) \wedge (\hat{a} > a) \wedge [R(a, \hat{a}) > Q(a, b)], \\ -\infty & \text{if } (a = 0) \wedge (b > -1), \end{cases} \end{aligned}$$

where $\wedge$ is the logical “and” and $\vee$ is the logical “or.” The value $+\infty$ comes from the fact that $u(z, a) \geq b$ does not constrain z when $a = 0$ and (24) does not constrain z when $\hat{a} = a$. Hence, the value of z can be driven to $+\infty$. The value $-\infty$ comes from two cases that we separate for clarity. In the first case, $z \leq Q(a, b)$ and $z \geq R(a, \hat{a})$ with $R(a, \hat{a}, b) > Q(a, b)$, leaving no choice for z and thus we set $z = -\infty$ to denote the maximizer of an empty set. In the second case, $a = 0$ and $b > 1$, so the constraint $u(z, a) \geq 0$ is assuredly violated and so again $z = -\infty$. The case where $z(a, \hat{a}, b) = R(a, \hat{a}, b)$ comes from the fact (23) does not constrain z when $a = 0$ as long as $u(z, 0) = -1 \geq b$. Since $\hat{a} < a$, z is driven to the upper bound $R(a, \hat{a}, b)$ from (24).

Step 2: Characterize actions. The next step is to solve

$$\inf_{\hat{a} \in [-2, 2]} v(z(a, \hat{a}, b), a).$$

As noted in Example 3, this infimum may not be attained and so we work with the possibility that no $\hat{a}(a, b)$ exists. For fixed a , $v(z(a, \hat{a}, b), a)$ is an increasing function of $z(a, \hat{a}, b)$ and so $\hat{a}$ should be chosen to minimize $z(a, \hat{a}, b)$. This eliminates the case where $z(a, \hat{a}, b) = +\infty$. A key step is to remove the dependence of $R(a, \hat{a}, b)$ on $\hat{a}$ through optimization. To this end, we define

$$R^\uparrow(a) \equiv \sup_{\hat{a} > a} R(a, \hat{a}),$$

$$R^\downarrow(a) \equiv \inf_{\hat{a} < a} R(a, \hat{a}).$$

Since $\hat{a}$ is chosen to minimize $z(a, \hat{a}, b)$, we have

$$z(a, b) \equiv \begin{cases} \min\{Q(a, b), R^\downarrow(a)\} & \text{if } (a \neq 0) \wedge [R^\uparrow(a) \leq Q(a, b)], \\ R^\downarrow(0) & \text{if } (a = 0) \wedge (b \leq -1), \\ -\infty & \text{if } (a = 0) \wedge (b > -1), \\ -\infty & \text{if } (a \neq 0) \wedge [R^\uparrow(a) > Q(a, b)]. \end{cases} \quad (25)$$

If it exists, we may set

$$\hat{a}(a, b) = \begin{cases} \hat{a}^\uparrow(a) & \text{if } (a \neq 0) \wedge [R^\uparrow(a) > Q(a, b)], \\ \hat{a}^\downarrow(a) & \text{otherwise,} \end{cases}$$

where

$$\hat{a}^\uparrow(a) \in \operatorname{argmax}_{\hat{a} > a} R(a, \hat{a}),$$

$$\hat{a}^\downarrow(a) \in \operatorname{argmin}_{\hat{a} < a} R(a, \hat{a})$$

if they exist. The rest of the development is not contingent on the existence of $\hat{a}(a, b)$, $\hat{a}^\uparrow(a)$, and $\hat{a}^\downarrow(a)$. In the case where the infimum is not attained, Lemma 4 can be used to determine $w(b)$ given $a(b)$ directly. Whether the infimum is attained or not depends

on b , but does not impact the analysis that follows, which simply works with the values $R^\uparrow(a)$ and $R^\downarrow(a)$.

Finally, we choose $a(b)$ to maximize $v(z(a, b), a)$. We first examine the choice of b . If b is such that $\inf_a(R^\uparrow(a) - Q(a, b)) > 0$, then $z(a, b) = -\infty$ and so $v(z(a, b), a)$ is $-\infty$, no matter the choice of a . Moreover, since $Q(a, b)$ is decreasing in b , any larger b will also not be chosen. Let $\bar{b} := \inf_{b \geq -2} \{\inf_a(R^\uparrow(a) - Q(a, b)) > 0\}$. As discussed, any $b > \bar{b}$ will not be chosen. To compute $\bar{b}$, we can use the expressions

$$R^\uparrow(a) = \begin{cases} 4a(1-a^2) & \text{if } 1/\sqrt{3} \leq a \leq 2, \\ \frac{4}{27}(9a - 5a^3 + \sqrt{2}(3-a^2)^{3/2}) & \text{if } 0 \leq a \leq 1/\sqrt{3}, \end{cases}$$

$$R^\downarrow(a) = \begin{cases} 4a(1-a^2) & \text{if } 1 \leq a \leq 2, \\ -\frac{4}{27}(9a - 5a^3 + \sqrt{2}(3-a^2)^{3/2}) & \text{if } 0 \leq a \leq 1. \end{cases}$$

The reader may verify that $\bar{b}$ is finite and strictly greater than 0. We can write an expression for $a(b)$ as

$$a(b) = \begin{cases} 0 & \text{if } -2 \leq b \leq -1, \\ a^\uparrow(b) & \text{if } -1 \leq b < \bar{b}, \\ \in [0, 2] & \text{if } b \geq \bar{b}, \end{cases} \quad (26)$$

where $a^\uparrow(b)$ is an optimal solution to

$$\begin{aligned} \max_{a \in (0, 2]} \min\{Q(a, b), R^\downarrow(a)\} a - 2a^2 \\ \text{s.t. } R^\uparrow(a) \leq Q(a, b). \end{aligned} \quad (27)$$

Our expression for $a(b)$ in (26) follows since if $b \leq -1$, then $v(z(a, b), a) < 0$ if $a > 0$ because we are in the first case of (25) and $\min\{Q(a, b), R^\downarrow(a)\} < 0$. Hence, $a(b) = 0$ since $v(z(a, b), a) = 0$. When $-1 \leq b < \bar{b}$, we cannot set $a = 0$; otherwise $z(a, b) = -\infty$ and the problem is infeasible. The only other option is the first case of (25) where $a(b)$ solves (27). Finally, when $b \geq \bar{b}$, then $z(a, b) = -\infty$ from (25) and so the choice of a is irrelevant.

With $a(b)$ as defined above we may write

$$z(b) \equiv z(a(b), b) = \begin{cases} R^\downarrow(0) & \text{if } -2 \leq b \leq -1, \\ \min\{Q(a^\uparrow(b), b), R^\downarrow(a^\uparrow(b))\} & \text{if } -1 \leq b < \bar{b}, \\ -\infty & \text{if } b \geq \bar{b} \end{cases}$$

and finally

$$\text{val(SAND}|b) = v(z(b), b) = \begin{cases} 0 & \text{if } -2 \leq b \leq -1, \\ z(b)a^\uparrow(b) - 2(a^\uparrow(b))^2 & \text{if } -1 \leq b < \bar{b}, \\ -\infty & \text{if } b \geq \bar{b}. \end{cases} \quad (28)$$

Since the original problem is feasible, we can eliminate $b \geq \bar{b}$ from consideration. In (28) we now have the first term in the “inner” minimization of (19) for determining b^* . The second term can be expressed as

$$\max_{a \in a^{\text{BR}}(z(b))} v(z(b), a). \quad (29)$$

We claim that $b = 0$ solves (19) in Step 3 of the sandwich procedure. To see this, we make the observation

$$b < 0 \text{ implies } a(b) < 1 \text{ and } z(b) < 0. \quad (30)$$

This follows by observing that when $b < 0$, there are two cases: $b \leq -1$ and $-1 < b < 0$. When $b \leq -1$, then $a(b) = 0$ and $z(b) = R^\downarrow(0) < 0$. When $-1 < b < 0$, observe that $\min\{Q(a, b), R^\downarrow(a)\} < 0$ for all $a \in (0, 2]$, and so $z(b) < 0$ and the objective function in (27) is decreasing in a , which implies that the constraint in (27) is tight; that is, $R^\uparrow(a) = Q(a, b)$. The reader may verify that this implies $a < 1$ and so $a(b) = a^\uparrow(b) < 1$. This yields (30).

Returning to (29), suppose $b < 0$. Consider the set $a^{\text{BR}}(z(b))$ when (from (30)) $z(b) < 0$. Taking the derivative of $u(z, a)$ with respect to a when $a \leq 1$ yields

$$\frac{\partial}{\partial a} u(z(b), a) = -z(b) - 4a(a^2 - 1) > 0$$

and so any $a \leq 1$ cannot be a best response to $z(b)$. This implies that $a(b)$ (which is greater than 1 from (30)) is not a best response to $z(b)$ and, hence,

$$\text{val}(\text{SAND}|b) > \max_{a \in a^{\text{BR}}(z(b))} v(z(b), a) \quad (31)$$

when $b < 0$. In Example 3, we showed $(\text{SAND}|b)$ when $b = 0$ is a tight relaxation. In particular, this means that $(z(0), a(0))$ is an optimal solution to (P) and thus $a(0)$ is a best response to $z(0)$. Thus,

$$\text{val}(\text{SAND}|0) = \max_{a \in a^{\text{BR}}(z(0))} v(z(0), a)$$

and so $b = 0$ is in the argmin in (19). Since (31) holds for any $b < 0$, this implies that $b^* = 0$. $\diamond$

5. NONEXISTENCE OF THE INNER MINIMIZATION AND THE RELATIONSHIP WITH THE FOA

In this section, we comment on a few connections between the sandwich approach and the FOA. We show how this relationship is connected to the issue of nonexistence of a minimizer to the inner minimization in the definition of $(\text{SAND}|b)$. We have already remarked (and Example 5 below confirms) that our procedure applies when the FOA is invalid. However, there is more to say about the connection between these two approaches.

The astute reader will have noticed that (SAND| b) does not include the first-order constraint (FOC(a)) common to both the FOA and Mirrlees's approach. The fact that (FOC(a)) is not present is connected to how we handled the agent's optimization problem via (4) and how this optimization was pulled into the objective in (Max-Max-Min| b). Indeed, the minimization over the alternate best response included in (Max-Max-Min| b) and (SAND| b) can be understood as our way to account for the optimality of the agent's best response. In this perspective, first-order conditions are not explicitly necessary in the formulation; they are implied when the sandwich approach is valid.

We already discussed the case when the inner minimization over $\hat{a}$ in (SAND| b) is not attained in Lemma 4, where the sandwich problem is equivalent to a problem with a local stationarity condition. In the case where the inner minimization is attained for some $\hat{a}^* \neq a^*$ and the first-best contract is not optimal (the remaining case), we recover first-order conditions via Lemma 7 when $\hat{a}^*$ is an interior point. In this case, $-\delta^*(a^*, \hat{a}^*, b)U_a(w(a^*, \hat{a}^*, b), \hat{a}^*) = 0$ and $\delta^*(a^*, \hat{a}^*, b) = 0$ would imply that the first-best contract is optimal, contradicting Assumption 3. We conclude that $U_a(w(a^*, \hat{a}^*, b), \hat{a}^*) = 0$. This implies that the first-order condition holds for $\hat{a}^*$. Since $U(w(a^*, \hat{a}^*, b), a^*) \geq U(w(a^*, \hat{a}^*, b), \hat{a}^*)$ from the no-jump constraint in (SAND| b), this further implies that $U_a(w(a^*, \hat{a}^*, b), a^*) = 0$ must also be satisfied since a^* will also be a best response (here we have assumed for simplicity that a^* is an interior point).

We examine this phenomenon from a more basic perspective. Suppose the sandwich approach is valid (for instance, because b is tight at optimality) and sandwich relaxation (SAND| b) has an optimal solution $(a^*, \hat{a}^*, w^*)$. Moreover, suppose (i) the Lagrangian multiplier $\delta(a^*, \hat{a}^*, b)$ from Lemma 3 is strictly positive and (ii) $\hat{a}^* < a^*$. Condition (ii) is reasonable since typically an alternate best response is to deviate to a lower effort level, not a higher effort level. Recall that cost is assumed to be nondecreasing in Assumption A1.8. In a special case we can show this formally.

PROPOSITION 3. *If the principal is risk-neutral and the FOA is not valid, then there exists an alternate best response $\hat{a}$ such that $\hat{a} < a^*$.*

In other words, with a risk-neutral principal, unless the FOA is valid, the agent will have a best-response “shirking” action. Observe that this assumption does not require any monotonicity assumptions on the output distribution f .

Given this scenario, we have the equivalence

$$\begin{aligned} \text{val}(\text{SAND}|b) &= \inf_{\hat{a} \in \mathbb{A}} \max_{w \geq \underline{w}} \{V(w, a^*) : U(w, a^*) \geq b, U(w, a^*) - U(w, \hat{a}) \geq 0\} \\ &= \inf_{\hat{a} \leq a^*} \max_{w \geq \underline{w}} \{V(w, a^*) : U(w, a^*) \geq b, U(w, a^*) - U(w, \hat{a}) \geq 0\}. \end{aligned}$$

To understand the above equivalence, we note that the $\leq$ direction is always true since the right-hand side has an additional restriction on the minimization, but $\hat{a} = \hat{a}^* \leq a^*$ attains the minimum that is achieved by the left-hand side problem.

The right-hand side problem above is equivalent to

$$\inf_{\hat{a} \leq a^*} \max_{w \geq \underline{w}} \left\{ V(w, a^*) : U(w, a^*) \geq b, \frac{U(w, a^*) - U(w, \hat{a})}{a^* - \hat{a}} \geq 0 \right\}$$

since $a^* - \hat{a} \geq 0$ in the range of choices for $\hat{a}$. Since $U(w, a)$ is differentiable in a , by the mean-value theorem, there exists an $\tilde{a} \in [\hat{a}, a^*]$ such that $(U(w, a^*) - U(w, \hat{a})) / (a^* - \hat{a}) = U_a(w, \tilde{a})$. Therefore, we have the equivalence

$$\begin{aligned}
 \text{val}(\text{SAND}|b) &= \inf_{\hat{a} \leq a^*} \max_{w \geq \underline{w}} \left\{ V(w, a^*) : U(w, a^*) \geq b, \frac{U(w, a^*) - U(w, \hat{a})}{a^* - \hat{a}} \geq 0 \right\} \\
 &= \max_{w \geq \underline{w}} \inf_{\hat{a} \leq a^*} \left\{ V(w, a^*) : U(w, a^*) \geq b, \frac{U(w, a^*) - U(w, \hat{a})}{a^* - \hat{a}} \geq 0 \right\} \\
 &= \max_{w \geq \underline{w}} \inf_{\tilde{a} \leq a^*} \{ V(w, a^*) : U(w, a^*) \geq b, U_a(w, \tilde{a}) \geq 0 \} \\
 &\leq \inf_{\tilde{a} \leq a^*} \max_{w \geq \underline{w}} \{ V(w, a^*) : U(w, a^*) \geq b, U_a(w, \tilde{a}) \geq 0 \} \\
 &\leq \max_{a \in \mathbb{A}} \inf_{\tilde{a} \leq a} \max_{w \geq \underline{w}} \{ V(w, a) : U(w, a) \geq b, U_a(w, \tilde{a}) \geq 0 \} \\
 &\leq \max_{a \in \mathbb{A}} \max_{w \geq \underline{w}} \{ V(w, a) : U(w, a) \geq b, U_a(w, a) \geq 0 \} \\
 &= \text{val}(\text{FOA}).
 \end{aligned} \tag{32}$$

The second equality follows from the tightness of b , the third equality uses the mean-value theorem, and the first inequality is simply the min-max inequality. Note that the constraint $U_a(w, \tilde{a}) \geq 0$ usually is binding for the problem $\max_{w \geq \underline{w}} \{ V(w, a^*) : U(w, a^*) \geq b, U_a(w, \tilde{a}) \geq 0 \}$, particularly if the principal is risk-neutral (Rogerson 1985, Jewitt 1988). Then

$$\begin{aligned}
 &\inf_{\tilde{a} \leq a^*} \max_{w \geq \underline{w}} \{ V(w, a^*) : U(w, a^*) \geq b, U_a(w, \tilde{a}) \geq 0 \} \\
 &= \inf_{\tilde{a} \leq a^*} \max_{w \geq \underline{w}} \{ V(w, a^*) : U(w, a^*) \geq b, U_a(w, \tilde{a}) = 0 \},
 \end{aligned}$$

which means that the sandwich relaxation must satisfy the stationary condition $U_a(w, \tilde{a}) = 0$ as a constraint. Note that in the FOA, $\tilde{a}$ must be taken as a^* and so is a weaker requirement.

Note that even when the sandwich approach is not valid, (32) reveals that it is a stronger relaxation than the FOA. Indeed, the FOA requires $U_a(w, a) = 0$, whereas the sandwich approach requires $U_a(w, \tilde{a}) = 0$, where $\tilde{a}$ is a minimizer. The latter is a more stringent condition to satisfy.

These observations provide an interpretation of the sandwich relaxation as a strengthening of the FOA, where we are required to satisfy an additional first-order condition over a worst-case choice of alternate best response.

There remains the question of how the sandwich procedure proceeds when the FOA is, in fact, valid. The next result shows that the two approaches are compatible in this case.

PROPOSITION 4. *When the FOA is valid, $\text{val}(\text{SAND}|\underline{U}) = \text{val}(\text{FOA}) = \text{val}(\underline{P})$. That is, both the sandwich approach and the FOA recover the optimal contract of the original problem.*

Observe that the validity of the FOA implies that the starting reservation utility $\underline{U}$ is tight at optimality. The next result reveals a partial converse in the case where the infimum in (SAND| b) is not attained. We emphasize that the MLRP assumption is needed to establish the following result, which we pull out of a proof of an earlier result that is stated and proved in the [Appendix](#).

PROPOSITION 5. *Suppose b is tight optimality and the sandwich problem (SAND| b) has solution (a^*, w^*) , where the inner minimization does not have a solution. Then, given the action a^* and with modified (IR) constraint $U(w, a^*) \geq b$, the FOA is valid. That is,*

$$\text{val}(P) = \max_{w \geq \underline{w}} \{V(w, a^*) : U(w, a^*) \geq b \text{ and } (\text{FOC}(a^*))\}$$

and the optimal solution to the right-hand side implements a^ .*

6. ADDITIONAL EXAMPLES

In this section, we provide three additional examples that further illustrate the sandwich procedure. The first example is where the FOA is invalid but nonetheless satisfies Assumptions 1–4 and so is amenable to the sandwich procedure.

EXAMPLE 5. Consider the following principal–agent problem. The distribution of output X is exponential with $f(x, a) = e^{-x/a}/a$ for $x \in \mathcal{X} = \mathbb{R}_+$ and $a \in \mathbb{A} := [1/10, 1/2]$. The principal is risk-neutral (and so $v(y) = y$), the value of output is $\pi(x) = x$, the agent's utility is $u(y) = 2\sqrt{y}$, the agent's cost of effort is $c(a) = 1 - (a - 1/2)^2$, and the outside reservation utility is $\underline{U} = 0$. The minimum wage is $\underline{w} = 1/16$. It is straightforward to check that Assumptions 1 and 2 are satisfied. Existence of an optimal solution is guaranteed by [Kadan et al. \(2017\)](#) and so [Assumption 3](#) is also satisfied. Finally, the monotonicity conditions in [Assumption 4](#) hold trivially for f . This means that [Theorems 1 and 2](#) apply.

Note also that the FOA is invalid. To see this, using the first-order condition $U_a(w, a) = 0$ to replace the original IC constraint, the resulting solution is $a^{\text{foa}} = 1/2$ and $w^{\text{foa}}(x) = 1/4$. Clearly, $w^{\text{foa}}(x)$ is a constant function and under $w^{\text{foa}}(x)$, the agent's optimal choice is $a = 1/10$, not $a^{\text{foa}} = 1/2$. Hence, the FOA is invalid.

Now we apply the sandwich procedure to derive an explicit solution. We start with $b = \underline{U}$ and show that [Proposition 1](#) holds.

Step 1. Characterize contract. According to [Lemma 3](#) the unique optimal contract to (SAND| $a, \hat{a}, \underline{U}$) is of the form

$$w_{\lambda, \delta}(a, \hat{a}, \underline{U}) = \left[\lambda + \delta \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) \right]^2,$$

assuming that $w(x) > \underline{w}$ for all x (we verify this is the case below). Plugging the above contract into the two constraints $U(w_{\lambda, \delta}(a, \hat{a}, \underline{U}), a) = \underline{U}$ and $U(w_{\lambda, \delta}(a, \hat{a}, \underline{U}), \hat{a}) = U(w_{\lambda, \delta}(a, \hat{a}, \underline{U}), \hat{a})$, we find

$$\lambda(a, \hat{a}, \underline{U}) = \frac{1}{2}(1 - (a - 1/2)^2),$$

$$\delta(a, \hat{a}, \underline{U}) = \frac{(2a - \hat{a})\hat{a}(a + \hat{a} - 1)}{2(a - \hat{a})^2}.$$

Step 2. Characterize actions. We plug $w_{\lambda(a, \hat{a}, \underline{U}), \delta(a, \hat{a}, \underline{U})}(a, \hat{a}, \underline{U})$ from [Step 1](#) into the principal's utility function to obtain the optimized Lagrangian from (20):

$$\mathcal{L}^*(a, \hat{a}|\underline{U}) = a - \frac{1}{4}[1 - (a - 1/2)^2]^2 - \frac{1}{4}(2a - \hat{a})\hat{a}(a + \hat{a} - 1)^2.$$

Now we solve the max-min problem in (21), where $\mathcal{L}^*(a, \hat{a}|\underline{U})$ is a fourth-order polynomial equation of $\hat{a}$ with first-order condition

$$\frac{\partial}{\partial \hat{a}} \mathcal{L}^*(a, \hat{a}|\underline{U}) = \frac{1}{4}(a + \hat{a} - 1)[\hat{a}(a + \hat{a} - 1) - (2a - \hat{a})(a + \hat{a} - 1) - 2(2a - \hat{a})\hat{a}] = 0.$$

This yields three solutions: $\hat{a} = a - 1$, $\hat{a} = (a + 1/2 - \sqrt{3a^2 - a + 1/4})/2$, and $\hat{a} = (a + \frac{1}{2} + \sqrt{3a^2 - a + 1/4})/2$. Since $\hat{a} \in [1/10, 1/2]$, the only feasible interior minimizer is

$$\hat{a}(a, \underline{U}) = \frac{1}{2} \left(a + \frac{1}{2} - \sqrt{3a^2 - a + 1/4} \right).$$

Plugging the $\hat{a}(a, \underline{U})$ into $\mathcal{L}^*$, we can solve the outer maximization problem in (21) over a , which yields $a^* = 1/2$ and, hence, $\hat{a}^* = (2 - \sqrt{2})/4$. This implies

$$w^*(x) = \left[\frac{1}{2} + \frac{1}{16} \left(1 - \frac{f\left(x, \frac{1}{4}(2 - \sqrt{2})\right)}{f(x, 1/2)} \right) \right]^2 = \left[\frac{1}{2} + \frac{1}{16} (1 - (2 + \sqrt{2})e^{-2x(1+\sqrt{2})}) \right]^2 > \frac{1}{16}.$$

Hence, [Proposition 1](#) implies that $\underline{U}$ is tight at optimality and so by [Theorem 1](#) we have found an optimal contract. $\diamond$

Second, the equivalence of the sandwich approach and the FOA when the FOA is valid (from [Proposition 4](#)) is illustrated by examining the classical example of [Hölmstrom \(1979\)](#).

EXAMPLE 6. The distribution of output X is exponential with $f(x, a) = e^{-x/a}/a$ for $x \in \mathcal{X} = \mathbb{R}_+$ and $a \in \mathbb{A} := [0, \bar{a}]$. The principal is risk-neutral (and so $v(y) = y$), the value of output is $\pi(x) = x$, the agent's utility is $u(y) = 2\sqrt{y}$, the agent's cost of effort is $c(a) = a^2$, minimum wage is $\underline{w} = 0$, and the outside reservation utility is $\underline{U} \geq 7^{-2/3}$.⁸

[Hölmstrom \(1979\)](#) showed that the FOA applies to this problem. Now we apply the sandwich procedure to derive an explicit solution.

Step 1. Characterize contract. According to [Lemma 3](#), the unique optimal contract to (SAND| $a, \hat{a}, b$) is of the form

$$w_{\lambda, \delta}(a, \hat{a}, \underline{U}) = \left[\lambda + \delta \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) \right]^2,$$

⁸This number is chosen to ensure that the minimum wage constraint is strictly satisfied at the optimal solution, as explicitly assumed in [Hölmstrom \(1979\)](#). For example, $\underline{U} = 0$ may lead to a positive probability that the contract equals $\underline{w}$.

assuming that $w(x) > \underline{w}$ for all x (we verify this is the case below). Plugging the above contract into the two constraints $U(w_{\lambda,\delta}(a, \hat{a}, \underline{U}), a) = \underline{U}$ and $U(w_{\lambda,\delta}(a, \hat{a}, \underline{U}), a) = U(w_{\lambda,\delta}(a, \hat{a}, \underline{U}), \hat{a})$ yields

$$\lambda(a, \hat{a}, \underline{U}) = \frac{1}{2}(a^2 + \underline{U}),$$

$$\delta(a, \hat{a}, \underline{U}) = \max\left\{0, \frac{(2a - \hat{a})\hat{a}(a^2 - \hat{a}^2)}{2(a - \hat{a})^2}\right\} = \max\left\{0, \frac{(2a - \hat{a})\hat{a}(a + \hat{a})}{2(a - \hat{a})}\right\}.$$

Step 2. Characterize actions. We plug $w_{\lambda(a, \hat{a}, \underline{U}), \delta(a, \hat{a}, \underline{U})}(a, \hat{a}, \underline{U})$ from **Step 1** into the principal's utility function to obtain the optimized Lagrangian from (20):

$$\mathcal{L}^*(a, \hat{a}|\underline{U}) = \begin{cases} a - \frac{1}{4}(a^2 + \underline{U})^2 - \frac{1}{4}(2a - \hat{a})\hat{a}(a + \hat{a})^2 & \text{if } \frac{(2a - \hat{a})\hat{a}(a + \hat{a})}{2(a - \hat{a})} > 0, \\ a - \frac{1}{4}(a^2 + \underline{U})^2 & \text{if } \frac{(2a - \hat{a})\hat{a}(a + \hat{a})}{2(a - \hat{a})} \leq 0. \end{cases}$$

Now we solve the max-min problem in (21), where $\mathcal{L}^*(a, \hat{a}|\underline{U})$ is a fourth-order polynomial equation of $\hat{a}$ with first-order condition

$$\frac{\partial}{\partial \hat{a}} \mathcal{L}^*(a, \hat{a}|\underline{U}) = -(a + \hat{a})(a^2 + 2a\hat{a} - 2\hat{a}^2) = 0.$$

This yields two solutions: $\hat{a} = (1 - \sqrt{3})a/2$ and $(1 + \sqrt{3})a/2$. Since $a > 0$, $\hat{a} = (1 - \sqrt{3})a/2$ is not feasible. It is not optimal to choose $\hat{a} \geq 2a$ as a minimizer, and so $-(2a - \hat{a})\hat{a}(a + \hat{a})^2/4 \geq 0$. Also $a \leq \hat{a} < 2a$ is not optimal, since with this choice, $\mathcal{L}^*(a, \hat{a}|\underline{U}) = a - (a^2 + \underline{U})^2/4$. So the minimizer should be taken in $0 \leq \hat{a} < a$, where $-(a + \hat{a})(a^2 + 2a\hat{a} - 2\hat{a}^2)$ is decreasing in $\hat{a}$. Therefore, the infimum is not attained and we have

$$\inf_{\hat{a}} \mathcal{L}^*(a, \hat{a}|\underline{U}) = a - \frac{1}{4}(a^2 + \underline{U})^2 - a^4.$$

This yields a solution $a^*(\underline{U})$ that is specified by the first-order condition of the above optimization problem:

$$1 - 5a^3 - 2a\underline{U} = 0,$$

where we may assume $\underline{U} \geq 7^{-2/3}$ so that

$$w^*(x=0) = \frac{1}{2}(a^{*2} + \underline{U}) - a^{*2} = \frac{1}{2}(\underline{U} - a^*(\underline{U})^2) \geq 0.$$

By l'Hôpital's rule, we have

$$\lim_{\hat{a} \rightarrow a} \delta(a, \hat{a}, \underline{U}) \left(1 - \frac{f(x, \hat{a})}{f(x, a)}\right) \rightarrow a^3 \left(\frac{x - a}{a^2}\right) = a(x - a),$$

so the optimal GMH contract according to the sandwich procedure is

$$w^*(x) = \frac{1}{2}a^{*2} + a^*(x - a^{*2}).$$

The resulting solution is consistent with the FOA solution, where the resulting Lagrangian multiplier for the first-order condition is $\mu(a) = a^3$ (Hölmstrom 1979) and the principal's value function is exactly the same:

$$V(w^{\text{foa}}(a), a) = a - \lambda(a)^2 - \mu(a)^2 \mathbb{E} \left(\frac{\partial \log f(X, a)}{\partial a} \right)^2 = a - \frac{1}{4}(a^2 + \underline{U})^2 - a^4.$$

This completes the example. $\diamond$

Note the similarity in the setups of Examples 5 and 6. The first can be seen as a relatively minor variation on the second, and yet the FOA approach fails in the first but holds in the second. In both cases the sandwich procedure applies. This illustrates, in a concrete way, aspects of the rigidity of the FOA and the robustness of the sandwich approach.

In our final example, we solve an adjustment of the problem proposed by Araujo and Moreira (2001), who show that the FOA fails but nonetheless construct an optimal solution by solving a nonlinear optimization problem with 20 constraints using Kuhn–Tucker conditions. Although this problem fails the conditions of Theorem 2 (it fails Assumption A1.1 since there are only two outcomes), we can nonetheless use our approach (specifically Lemma 2 and Proposition 2) to construct an optimal contract. We remark that this example has the nice feature that all best responses are interior to the interval of actions $\mathbb{A} = [-1, 1]$, in contrast to all previous examples. Moreover, there are multiple alternate best responses. As can be seen below, and in relation to remarks in Section 5, stationarity conditions at these interior points are implicitly recovered via the sandwich approach.

EXAMPLE 7. The principal has expected utility $V(w, a) = \sum_{i=1}^2 p_i(a)(x_i - w_i)$, where $p_1(a) = a^2$, $p_2(a) = 1 - a^2$ for $a \in [-1, 1]$, and there are two possible outcomes $x_1 = 1$ and $x_2 = 3/4$, denoting $w_i = w(x_i)$ for $i = 1, 2$. The minimum wage is $\underline{w} = 0$. The agent's expected utility is $U(w, a) = \sum_{i=1}^2 p_i(a)\sqrt{w_i} - 2a^2(1 - 2a^2 + (4/3)a^4)$ with reservation utility $\underline{U} = 0$. We apply Step 1 and Step 2 of the sandwich procedure.

Step 1. Characterize contract. The first-order conditions (8) imply that an optimal solution (SAND| $a, \hat{a}, 0$) must satisfy

$$w_i^* = w^*(x_i) = \frac{1}{4} \left[\lambda + \delta \left(1 - \frac{p_i(\hat{a})}{p_i(a)} \right) \right]^2 \quad \text{for } i = 1, 2, \quad (33)$$

assuming that $w_i^* \geq \underline{w}$ for $i = 1, 2$ (we check below that this is the case) for some choice of λ and δ . To characterize these λ and δ , we plug the above contract into the two constraints of (SAND| $a, \hat{a}, b$), $U(w^*, a) = 0$ and $U(w^*, a) = U(w^*, \hat{a})$, to find

$$\begin{aligned} \lambda(a, \hat{a}, 0) &= 4a^2 \left(1 - 2a^2 + \frac{4}{3}a^4 \right), \\ \delta(a, \hat{a}, 0) &= \frac{4a^2(1 - a^2)[3 + 4a^4 + 4\hat{a}^4 + 4a^2\hat{a}^2 - 6(\hat{a}^2 + a^2)]}{3(a^2 - \hat{a}^2)}. \end{aligned} \quad (34)$$

Step 2. Characterize actions. We solve (21), where

$$\begin{aligned}
 \mathcal{L}^*(a, \hat{a}|0) &= \sum_{i=1}^2 p_i(a)(x_i - w(a, \hat{a}, 0)_i) \\
 &= \sum_{i=1}^2 p_i(a)x_i - \frac{1}{4}\lambda(a, \hat{a}, 0)^2 - \frac{1}{4} \frac{\delta(a, \hat{a}, 0)^2}{\sum_{i=1}^2 \left(1 - \frac{p_i(\hat{a})}{p_i(a)}\right)^2 p_i(a)} \\
 &= a^2 + \frac{3}{4}(1 - a^2) - \frac{4}{9}a^4(3 - 6a^2 + 4a^4)^2 \\
 &\quad - \frac{4}{9}a^2(1 - a^2)[3 + 4a^4 + 4\hat{a}^4 + 4a^2\hat{a}^2 - 6(\hat{a}^2 + a^2)]^2
 \end{aligned}$$

by leveraging Lemma 7. Note that only the last term $t(a, \hat{a}) \equiv [3 + 4a^4 + 4\hat{a}^4 + 4a^2\hat{a}^2 - 6(\hat{a}^2 + a^2)]^2$ in the last line of the above expression involves $\hat{a}$. By taking the first-order condition with respect to $\hat{a}$, we obtain three solutions:

$$\hat{a} = 0, \quad \hat{a} = \frac{\sqrt{3 - 2a^2}}{2}, \quad \hat{a} = -\frac{\sqrt{3 - 2a^2}}{2}.$$

One can verify that for any $a \in [-1, 1]$,

$$t(a, 0) = (3 - 6a^2 + 4a^4)^2 < \frac{9}{16}(1 - 2a^2)^4 = t\left(a, \frac{\sqrt{3 - 2a^2}}{2}\right) = t\left(a, -\frac{\sqrt{3 - 2a^2}}{2}\right).$$

Therefore, the unique minimizer of $\mathcal{L}^*(a, \hat{a}|0)$ over $\hat{a}$ is $\hat{a}^*(a) \equiv 0$. Then

$$\mathcal{L}^*(a, 0|0) = a^2 + \frac{3}{4}(1 - a^2) - \frac{4}{9}a^4(3 - 6a^2 + 4a^4)^2 - \frac{4}{9}a^2(1 - a^2)[3 + 4a^4 - 6a^2]^2$$

has a maximum at $a^* = \sqrt{3}/2$ (there are three maximizers, $a^* = \pm\sqrt{3}/2$ and $a^* = 0$, all interior to $\mathbb{A}$; we just pick $a^* = \sqrt{3}/2$). This completes the sandwich procedure and we have produced an optimal solution to (SAND|0) of the form $(a^*, \hat{a}^*, w^*)$, where $a^* = \sqrt{3}/2$, $\hat{a}^* = 0$, and $w_1^* = 1$ and $w_2^* = 0$ (using the fact $\lambda(\sqrt{3}/2, 0, 0) = 3/4$ and $\delta(\sqrt{3}/2, 0, 0) = 1/4$). Note, in particular, that $w_i^* \geq \underline{w} = 0$ for $i = 1, 2$.

Second, we show that (w^*, a^*) is feasible to (P). It suffices to show that a^* is a best response to w^* . The agent's expected utility under the contract $w^* = w(a, \hat{a}, 0)$ and taking action $\tilde{a}$ is (using (33) and (34))

$$U(w^*, \tilde{a}) = \frac{4}{3}(a^2 - \tilde{a}^2)(\tilde{a}^2 - \hat{a}^2)(2a^2 + 2\hat{a}^2 + 2\tilde{a}^2 - 3).$$

Given $a^* = \sqrt{3}/2$ and $\hat{a}^* = 0$, $U(w^*, \tilde{a})$ is indeed maximized at $\tilde{a} = \pm\sqrt{3}/2$ and $\tilde{a} = 0$. This shows that a^* is a best response to w^* and, hence, (w^*, a^*) is feasible to (P).

Finally, by Lemma 2 we know that $\text{val}(\text{SAND}|0) \geq \text{val}(\text{P})$ and this implies that (w^*, a^*) achieves the best possible principal utility in (P). We conclude that w^* is an optimal

contract. However, one can check that the FOA is not valid. The solution to (FOA) will yield $a^{\text{foa}} = 0.798$, which cannot be implemented by the corresponding w^{foa} . Details are suppressed. $\diamond$

7. CONCLUSION

We provide a general method to solve moral hazard problems when output is a continuous random variable with a distribution that satisfies the MLRP (Assumption 4). This involves solving a tractable relaxation of the original problem using a bound on agent utility derived from our proposed procedure.

We do admit that, in general, Step 3 of the sandwich procedure may be a priori difficult unless sufficient structural information is known about the set $a^{\text{BR}}(w(b))$. However, as the examples in this paper illustrate, this may not be an issue in sufficiently well behaved cases.

Indeed, Proposition 1 is helpful in this regard. If one initially assumes $b = \underline{U}$, and the resulting optimal contract can be shown to strictly satisfy the limited liability constraint, there is no need for Step 3 of the sandwich procedure. This method is on display in Example 5. As mentioned in the paper, there is precedence in the moral hazard literature to restrict analysis to cases where the limited liability constraint is guaranteed not to bind.

Finding additional criteria for the (IR) constraint to be tight is an important area for further investigation. Moreover, finding other scenarios where (19) is tractable is also of interest. Examples 4–7 show that the basic framework of our approach can help solve problems that may not satisfy all the assumptions used in our theorems.

APPENDIX: PROOFS

A.1 Proof of Lemma 1

We set the notation $V^*(a, \hat{a}) = \max_{w \geq \underline{w}} \{V(w, a) : (w, a) \in \mathcal{W}(\hat{a}, b)\}$ and $V^*(a) = \inf_{\hat{a}} V^*(a, \hat{a})$. The result follows by establishing the following claim.

CLAIM 1. *The term $V^*(a)$ is upper semicontinuous in a .*

Indeed, if $V^*(a)$ is upper semicontinuous, then, since $\mathbb{A}$ is compact, an outer maximizer a exists by the Weierstrass theorem.

We now establish the claim. By definition of upper semicontinuity, we want to show that for any constant $\alpha \in \mathbb{R}$, $\{a | V^*(a) < \alpha\}$ is open, where α is independent of a . This shows that there exists an $\epsilon > 0$ such that for all $a' \in \mathcal{N}_\epsilon(a)$, $V^*(a') < \alpha$, where $\mathcal{N}_\epsilon(a)$ is an open neighborhood of a . Now we pick any $a_0 \in \{a | V^*(a) < \alpha\}$. Note that $\inf_{\hat{a}} V(a_0, \hat{a}) < \alpha$ implies that there exists some $\hat{a}_0$ such that

$$V(a_0, \hat{a}_0) < \alpha.$$

Alternatively, since $V(a, \hat{a})$ is upper semicontinuous by the theorem of maximum, the set

$$\{(a, \hat{a}) | V(a, \hat{a}) < \alpha\}$$

is open. Therefore, there exists an $\epsilon > 0$ such that $V(a', \hat{a}') < \alpha$ for any $(a', \hat{a}') \in \mathcal{B}_\epsilon(a_0, \hat{a}_0)$, where $\mathcal{B}_\epsilon(a_0, \hat{a}_0)$ is the open ball in $\mathbb{R}^2$ centered at $(a_0, \hat{a}_0)$ with radius ϵ . Thus, we can find an open neighborhood $\mathcal{N}_{\epsilon_1}(a_0)$ of a_0 and $\mathcal{N}_{\epsilon_2}(\hat{a}_0)$ of $\hat{a}_0$ such that

$$\mathcal{N}_{\epsilon_1}(a_0) \times \mathcal{N}_{\epsilon_2}(\hat{a}_0) \subseteq \mathcal{B}_\epsilon(a_0, \hat{a}_0).$$

Therefore, we have $V(a', \hat{a}') < \alpha$ for any $a' \in \mathcal{N}_{\epsilon_1}(a_0)$ and $\hat{a}' \in \mathcal{N}_{\epsilon_2}(\hat{a}_0)$. As a result, for any $a' \in \mathcal{N}_{\epsilon_1}(a_0)$, we have

$$V^*(a') = \inf_{\hat{a}} V(a', \hat{a}) \leq V(a', \hat{a}') < \alpha$$

for a given $\hat{a}' \in \mathcal{N}_{\epsilon_2}(\hat{a}_0)$, which shows that $\{a | V^*(a) < \alpha\}$ is open. We thus obtain the desired upper semicontinuity of $\inf_{\hat{a}} V(a, \hat{a})$.

A.2 Proof of Lemma 2

Observe that

$$\begin{aligned} \text{val}(\mathbf{P}|b) &= \text{val}(\text{Max-Max-Min}|b) \\ &= \max_{a \in \mathbb{A}} \max_{w \geq \underline{w}} \inf_{\hat{a} \in \mathbb{A}} V^I(w, a|\hat{a}, b) \\ &\leq \max_{a \in \mathbb{A}} \inf_{\hat{a} \in \mathbb{A}} \max_{w \geq \underline{w}} V^I(w, a|\hat{a}, b) \\ &= \text{val}(\text{SAND}|b), \end{aligned}$$

where the inequality follows by the min-max inequality. If there exists an optimal solution (w^*, a^*) to $(\mathbf{P})$ such that $U(w^*, a^*) \geq b$ (and thus is also a feasible solution to $(\mathbf{P}|b)$), then $\text{val}(\mathbf{P}) \leq \text{val}(\mathbf{P}|b)$. However, we already argued in the main text that $\text{val}(\mathbf{P}) \geq \text{val}(\mathbf{P}|b)$. This implies $\text{val}(\mathbf{P}) = \text{val}(\mathbf{P}|b)$ and so the above inequality implies $\text{val}(\mathbf{P}) \leq \text{val}(\text{SAND}|b)$.

A.3 Proof of Lemma 3

The proof of (i) and (ii) is analogous to the proof of Theorem 3 in Ke and Ryan (2018). In both cases, a , $\hat{a}$, and b are fixed constants. The difference here is that the no-jump constraint defining $(\text{SAND}|b)$ is an inequality, while in Ke and Ryan (2018) the no-jump constraint is an equality. However, this only changes the complementary slackness properties, as detailed in the lemma. Moreover, in Ke and Ryan (2018) we need not entertain the case where $\hat{a} = a$. Fortunately, the case where $\hat{a} = a$ is straightforward, since then $(\text{SAND}|a, \hat{a}, b)$ is solved by the first-best contract, which is unique. Further details are omitted.

The proof of (iii) and (iv) is standard by applying the theorem of maximum. Details are omitted.

Assumption 2 is required in the proof of Theorem 3 in Ke and Ryan (2018), and that is also why it is required here.

A.4 Proof of [Lemma 4](#)

For convenience, we denote

$$V^*(a, \hat{a}|b) = \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U(w, a) - U(w, \hat{a}) \geq 0\}.$$

If $\inf_{\hat{a}} V^*(a, \hat{a}|b)$ is not attained, it must be that the infimizing sequence converges to a (for more details on this argument, see the discussion following [Lemma 3](#) is the main text of the paper). We can decompose the minimization problem as

$$\begin{aligned} & \inf_{\hat{a}} \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U(w, a) - U(w, \hat{a}) \geq 0\} \\ &= \inf \left\{ \inf_{\hat{a} \leq a} V^*(a, \hat{a}|b), \inf_{\hat{a} \geq a} V^*(a, \hat{a}|b) \right\}. \end{aligned}$$

Case 1: $\inf_{\hat{a} \leq a} V^*(a, \hat{a}|b) = \inf_{\hat{a}} V^*(a, \hat{a}|b)$. We begin by observing that if $\inf_{\hat{a} \leq a} V^*(a, \hat{a}|b)$ has an infimizing sequence that does not converge to a , then by the supposition of nonexistence, we must have

$$\inf_{\hat{a} \leq a} V^*(a, \hat{a}|b) > \inf_{\hat{a}} V^*(a, \hat{a}|b).$$

In this case, we switch to consider $\inf_{\hat{a} \geq a} V^*(a, \hat{a}|b)$, which is discussed in Case 2 below.

By the mean-value theorem, there exists an $\tilde{a} \in [\hat{a}, a]$ such that $(U(w, a) - U(w, \hat{a})) / (a - \hat{a}) = U_a(w, \tilde{a})$. Therefore, we have the equivalence

$$\begin{aligned} \inf_{\hat{a} \leq a} V^*(a, \hat{a}|b) &= \inf_{\hat{a} \leq a} \max_{w \geq \underline{w}} \left\{ V(w, a^*) : U(w, a) \geq b, \frac{U(w, a) - U(w, \hat{a})}{a - \hat{a}} \geq 0 \right\} \\ &= \lim_{\hat{a} \rightarrow a^-} \max_{w \geq \underline{w}} \left\{ V(w, a^*) : U(w, a) \geq b, \frac{U(w, a) - U(w, \hat{a})}{a - \hat{a}} \geq 0 \right\} \quad (35) \\ &= \lim_{\tilde{a} \rightarrow a^-} \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U_a(w, \tilde{a}) \geq 0\}. \end{aligned}$$

Note that $\max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U_a(w, \tilde{a}) \geq 0\}$ is continuous in $\tilde{a}$ (since U is continuously differentiable in a) and, as mentioned above, the infimizing sequence converges to a and so a minimizer exists to (35), yielding

$$\inf_{\hat{a} \leq a} V^*(a, \hat{a}|b) = \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U_a(w, a) \geq 0\}.$$

It remains to show that the constraint $U_a(w, a) \geq 0$ is binding for any $a \in \text{int } \mathbb{A}$ and slack only if $a = \tilde{a}$. Suppose that the constraint in the above problem is slack at the optimum, i.e., $U_a(w, a) > 0$. Then the Lagrangian multiplier is zero and we have

$$\max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U_a(w, a) \geq 0\} = \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b\}.$$

This means that $w^{\text{fb}}(a|b)$ solves $\max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U_a(w, a) \geq 0\}$, where $w^{\text{fb}}(a|b)$ is the first-best contract. Equivalently, we have

$$\inf_{\hat{a}} V^*(a, \hat{a}|b) = \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b\}. \quad (36)$$

We now claim that $w^{\text{fb}}(a|b)$ implements a . Continuing from (36) and letting $\hat{a}' \in a^{\text{BR}}(w^{\text{fb}}(a|b))$, we have

$$\inf_{\hat{a}} V^*(a, \hat{a}|b) \leq V^*(a, \hat{a}'|b) \leq \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b\} = \inf_{\hat{a}} V^*(a, \hat{a}|b),$$

where the first inequality is by the definition of minimization and the second inequality is straightforward by withdrawing a constraint from the maximization problem. Therefore, all inequalities become equalities and $w^{\text{fb}}(a|b)$ satisfies the no-jump constraint $U(w, a) - U(w, \hat{a}') \geq 0$. This implies $a \in a^{\text{BR}}(w^{\text{fb}}(a|b))$. Therefore, for any $a \in \text{int } \mathbb{A}$, $U_a(w^{\text{fb}}(a|b), a) = 0$, and $U_a(w^{\text{fb}}(a|b), a) > 0$ only occurs when $a = \bar{a}$, where $w^{\text{fb}}(\bar{a}|b)$ implements $\bar{a}$. This completes Case 1.

Case 2: $\inf_{\hat{a} \geq a} V^*(a, \hat{a}|b) = \inf_{\hat{a}} V^*(a, \hat{a}|b)$. In this case, we have the equivalence

$$\begin{aligned} \inf_{\hat{a} \geq a} V^*(a, \hat{a}|b) &= \lim_{\hat{a} \rightarrow a^+} \max_{w \geq \underline{w}} \left\{ V(w, a^*) : U(w, a) \geq b, \frac{U(w, a) - U(w, \hat{a})}{a - \hat{a}} \leq 0 \right\} \\ &= \lim_{\hat{a} \rightarrow a^+} \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b, U_a(w, \tilde{a}) \leq 0\}. \end{aligned}$$

The rest of the argument is quite similar to Case 1 and thus is omitted. Combining these two cases, we have the desired conclusion.

A.5 Proof of Lemma 5

We require the following lemma.

LEMMA 9 (Theorem 6 in Section 8.5 of Lasdon (2011)). *Consider a maximization problem*

$$\max_y \{f(y) : g(y) \geq 0\},$$

where $f : \mathbb{Y} \rightarrow \mathbb{R}$ and $g : \mathbb{Y} \rightarrow \mathbb{R}^k$ for some compact subset $\mathbb{Y} \subset \mathbb{R}^n$. Assume that both f and g are continuous and differentiable. If the Lagrangian $L(y, \alpha) = f(y) + \alpha \cdot g(y)$ is strictly concave in y , then

$$\max_y \{f(y) : g(y) \geq 0\} = \inf_{\alpha \geq 0} \max_y L(y, \alpha),$$

where we assume the maximum of $L(y, \alpha)$ over y exists for any given α .

PROOF OF LEMMA 5. When the infimum in (SAND|b) is not attained or attained at $a^\#$, the result follows a standard application of duality theory via Lemma 9, due to Lemma 4.

We now consider the case where the infimum is attained. Let $(a^*, \hat{a}^*, z^*)$ be an optimal solution (SAND|b); that is,

$$V(z^*, a^*) = \max_{a \in \mathbb{A}} \inf_{\hat{a} \in \mathbb{A}} \max_{z \geq \underline{z}} \{V(z, a) : U(z, a) \geq b, U(z, a) - U(z, \hat{a}) \geq 0\}.$$

Given a^* , consider the Lagrangian dual of the inner maximization problem over z ; that is,

$$\mathcal{L}(z, \lambda, \delta|a^*, \hat{a}^*, b) = V(z, a^*) + \lambda[U(z, a^*) - b] + \delta[U(z, a^*) - U(z, \hat{a}^*)].$$

Note that $\mathcal{L}$ is strictly concave in z since $V(z, a^*) = v(\pi(x_0) - z)$ is concave and $U(z, a^*) = u(z)$ is strictly concave in z , and the term involving δ is a function only of a since $U(z, a^*) - U(z, \hat{a}) = u(z) - c(a^*) - (u(z) - c(\hat{a})) = c(\hat{a}) - c(a^*)$. Lemma 9 implies

$$\begin{aligned} & \inf_{\hat{a} \in \mathbb{A}} \max_{z \geq \underline{z}} \{V(z, a^*) : U(z, a^*) \geq d, U(z, a^*) - U(z, \hat{a}) \geq 0\} \\ &= \inf_{\hat{a} \in \mathbb{A}} \inf_{\lambda, \delta \geq 0} \max_{z \geq \underline{z}} \mathcal{L}(z, \lambda, \delta|a^*, \hat{a}, d) \end{aligned} \quad (37)$$

for all $d \in [b, b + \epsilon)$. We now consider three cases. We show that the first two cases do not occur, leaving only the third case where we can establish the result. The cases consider how perturbing b can effect the primal and dual problems in (37).

Case 1. The set $\bigcap_{\hat{a} \in \mathbb{A}} \{z : U(z, a^*) \geq b + \epsilon, U(z, a^*) - U(z, \hat{a}) \geq 0\}$ is empty for any arbitrarily small $\epsilon > 0$. Here, the Lagrangian multiplier

$$\lambda(a^*, \hat{a}_\epsilon^*) \in \arg \inf_{\lambda, \delta \geq 0} \max_{z \geq \underline{z}} \mathcal{L}(z, \lambda, \delta|a^*, \hat{a}, b + \epsilon)$$

is unbounded, where $\hat{a}_\epsilon^* \in \arg \min_{\hat{a}} \inf_{\lambda \geq 0, \delta \geq 0} \max_{z \geq \underline{z}} L(z, \lambda, \delta|a^*, \hat{a}, b + \epsilon)$. Also, $U(z_\epsilon^*, a^*) < b + \epsilon$ for any z_ϵ^* such that

$$\mathcal{L}(z_\epsilon^*, \lambda(a^*, \hat{a}_\epsilon^*), \delta(a^*, \hat{a}_\epsilon^*)|a^*, \hat{a}_\epsilon^*, b + \epsilon) = \inf_{\lambda \geq 0} \inf_{\delta \geq 0} \max_{z \geq \underline{z}} \mathcal{L}(z, \lambda, \delta|a^*, \hat{a}, b + \epsilon).$$

Therefore, we choose a sequence $\epsilon_n = \epsilon/n$ and we have

$$U(z_{\epsilon_n}^*, a^*) - b - \epsilon_n < 0,$$

where $z_{\epsilon_n}^*$ is a sequence such that

$$V(z_{\epsilon_n}^*, a^*) = \inf_{\hat{a} \in \mathbb{A}} \inf_{\lambda \geq 0, \delta \geq 0} \max_{z \geq \underline{z}} L(z, \lambda, \delta|a^*, \hat{a}, b + \epsilon_n).$$

Note that $(z_\epsilon^*, a^*, \hat{a}_\epsilon^*)$ is upper hemicontinuous in ϵ by the theorem of maximum. Then as $n \rightarrow \infty$, the limit $(z_0^*, a_0^*, \hat{a}_0^*; \lambda(a_0^*, \hat{a}_0^*), \delta(a_0^*, \hat{a}_0^*))$ is a solution to the problem without perturbation ($\epsilon = 0$). Without loss of generality, we choose

$$(z^*, a^*, \hat{a}^*; \lambda(a^*, \hat{a}^*), \delta(a^*, \hat{a}^*)) = (z_0^*, a_0^*, \hat{a}_0^*; \lambda(a_0^*, \hat{a}_0^*), \delta(a_0^*, \hat{a}_0^*)).$$

Then, passing to the limit (taking a subsequence if necessary), $z_{\epsilon_n}^* \rightarrow z^*$ and we have

$$\lim_{n \rightarrow \infty} [U(z_{\epsilon_n}^*, a^*) - b - \epsilon_n] = U(z^*, a^*) - b \leq 0,$$

which contradicts the supposition $U(z^*, a^*) > b$.

Case 2. The set $\bigcap_{\hat{a} \in \mathbb{A}} \{z : U(z, a^*) \geq b + \epsilon, U(z, a^*) - U(z, \hat{a}) \geq 0\}$ is nonempty and $\lambda(a^*, \hat{a}_\epsilon^*) > 0$, for any $\epsilon > 0$. Note that $\lambda(a^*, \hat{a}_\epsilon^*) > 0$ implies that the constraint $U(z_\epsilon^*, a^*) \geq$

$b + \epsilon$ is binding given strong duality. We choose a sequence $\epsilon_n = \epsilon/n$. Passing to the limit (taking a subsequence if necessary), $z_{\epsilon_n}^* \rightarrow z^*$ and we have

$$0 = \lim_{n \rightarrow \infty} [U(z_{\epsilon_n}^*, a^*) - b - \epsilon_n] = U(z^*, a^*) - b,$$

which contradicts with the supposition $U(z^*, a^*) > b$.

Case 3. The set $\bigcap_{\hat{a} \in A} \{z : U(z, a^*) \geq U^* + \epsilon, U(z, a^*) - U(z, \hat{a}) \geq 0\}$ is nonempty and $\lambda(a^*, \hat{a}_\epsilon^*) = 0$ for some arbitrarily small $\epsilon > 0$. Given $\lambda(a^*, \hat{a}_\epsilon^*) = 0$, we have

$$\begin{aligned} V(z_\epsilon^*, a^*) &= \max_z V(z, a^*) + \lambda(a^*, \hat{a}_\epsilon^*)(U(z, a^*) - b - \epsilon) + \delta(a^*, \hat{a}_\epsilon^*)(U(z, a^*) - U(z, \hat{a}_\epsilon^*)) \\ &= \max_z V(z, a^*) + \lambda(a^*, \hat{a}_\epsilon^*)(U(z, a^*) - b) + \delta(a^*, \hat{a}_\epsilon^*)(U(z, a^*) - u(z, \hat{a}_\epsilon^*)) \\ &\geq \inf_{\hat{a}} \inf_{\lambda, \delta \geq 0} \max_z V(z, a^*) + \lambda(U(z, a^*) - b) + \delta(U(z, a^*) - U(z, \hat{a}_\epsilon^*)) \\ &= V(z^*, a^*). \end{aligned}$$

We already know that $V(z^*, a^*) \geq V(z_\epsilon^*, a^*)$ since $\epsilon > 0$. Therefore, we have shown $V(z_\epsilon^*, a^*) = V(z^*, a^*)$, as required.

The above argument shows that as we increase b to $b + \epsilon$, we can find a new optimal contract that does not change the objective value. This can be repeated until we find a sufficiently large ϵ such that $U(z_\epsilon^*, a_\epsilon^*) = b + \epsilon$. $\square$

A.6 Proof of *Theorem 1*

There are two cases to consider. The first is when the inner inf in (SAND| b) is not attained. This is handled by the following lemma.

LEMMA 10. *Suppose b is tight at optimality and the sandwich problem (SAND| b) has solution (a^*, w^*) , where the inner minimization does not have a solution. Then, given the action a^* and with modified (IR) constraint $U(w, a^*) \geq b$, the FOA is valid. That is,*

$$\text{val}(P) = \max_{w \geq \underline{w}} \{V(w, a^*) : U(w, a^*) \geq b \text{ and } (\text{FOC}(a^*))\} = \text{val}(\text{SAND}|b).$$

PROOF. We first argue that $a^{\text{BR}}(w(b))$ is not a singleton. Suppose there exists an $\hat{a}^* \neq a^*$ such that the GMH contract $w(a^*, \hat{a}^*, b)$ implements a^* (see Proposition 6 and also Remark 4 in Ke and Ryan 2018), i.e., $V(w(a^*, \hat{a}^*, b), a^*) = \text{val}(P)$. Note that for any $\hat{a} \in \mathbb{A}$,

$$\text{val}(\text{SAND}|a^*, \hat{a}, b) \geq \max_{(w, a^*)} \{V(w, a^*) : U(w, a^*) \geq U^*, a^* \in a^{\text{BR}}(w)\}.$$

Therefore, $\hat{a}^*$ is the solution to the inner minimization problem

$$\hat{a}^* \in \arg \min_{\hat{a}} V^*(a^*, \hat{a}|b),$$

which contradicts the supposition of nonexistence. Therefore, the best-response set $a^{\text{BR}}(w(b))$ must be a singleton, i.e., a^* is the unique best response at the optimal solution. In this case, according to Mirrlees (1999), all no-jump constraints are slack at optimality and the FOA is valid (up to the modified IR constraint $U(w, a^*) \geq b$).

Finally, by [Lemma 4](#), we know that $\text{val}(\text{SAND}|b)$ is equal to the value of the FOA with modified IR constraint $U(w, a^*) \geq b$. $\square$

We now return to the case where the infimum in $(\text{SAND}|b)$ is attained. The proof proceeds in two stages. In the first stage, we examine a subclass of problems where the agent's action a is given. In the second stage, we illustrate how to determine the right choice for a .

REMARK 1. We remark that the analysis of the first stage of the proof is drawn from results in [Ke and Ryan \(2018\)](#). In that paper it is assumed that an action a^* is given and is implemented by an optimal contract w^* such that $U(w^*, a^*) = \underline{U}$. In this setting, the assumption that $U(w^*, a^*) = \underline{U}$ is without loss of interest, since we assume that a^* and w^* are given, and so $\underline{U}$ can be defined as $U(w^*, a^*)$. The assumption that $U(w^*, a^*) = \underline{U}$ is critical in Section 4 of [Ke and Ryan \(2018\)](#). See Remark 4 of that paper for further discussion of this point. This is an important difference with our current analysis. Here we no longer assume that a target a^* is given and so we cannot assume without loss of generality that $U(w^*, a^*) = \underline{U}$. Indeed, uncovering a method to *find* w^* and a^* is the focus of this paper.

Accordingly, the analysis here proceeds in a different manner than [Ke and Ryan \(2018\)](#). First, [Ke and Ryan \(2018\)](#) consider a simpler version of $(\text{Min-Max}|a, b')$ where the no-jump constraint is an equality. This is sufficient in that setting because they do not need to further analyze this problem to determine a^* ; it is simply given. This oversimplifies the current development. Moreover, Stage 2 is not needed to analyze the situation in [Ke and Ryan \(2018\)](#). The added complexity of Stage 2 arises precisely because the optimal action for the agent and the utility delivered to the agent at optimality are both a priori unknown.

A.6.1 Analysis of Stage 1 Define the following intermediate problem, which is the parametric problem $(P|b')$ with $b' \geq \underline{U}$ and the agent's action a fixed:

$$\begin{aligned} & \max_{w \geq \underline{w}} V(w, a) \\ & \text{subject to } U(w, a) \geq b' \quad (P|a, b') \\ & U(w, a) - U(w, \hat{a}) \geq 0 \quad \text{for all } \hat{a} \in \mathbb{A}. \end{aligned}$$

We restrict attention to where the above problem is feasible; that is, a is an implementable action that delivers at least utility b' to the agent. Note that we need not take b' equal to the b that is tight at optimality provided in the hypothesis of the theorem. It is an arbitrary $b' \geq \underline{U}$ with the above property.

We can define the related problem

$$\inf_{\hat{a} \in \mathbb{A}} \max_{w \geq \underline{w}} \{V(w, a) : U(w, a) \geq b', U(w, a) - U(w, \hat{a}) \geq 0\}. \quad (\text{Min-Max}|a, b')$$

We denote an optimal solution to $(\text{Min-Max}|a, b')$ by $\hat{a}(a, b')$ and $w_{a, b'}$.

Note that $(P|a, b')$ is analogous to $(P|b)$ and $(\text{Min-Max}|a, b')$ is analogous to $(\text{SAND}|b)$, however with a given. The key result is an implication of Theorem 4 in [Ke and Ryan \(2018\)](#) carefully adapted to this setting.

PROPOSITION 6. *Suppose Assumptions 1–4 hold. Let a be an implementable action and let $b' = U(w^{a, \underline{U}}, a)$, where $w^{a, \underline{U}}$ is an optimal solution to $(P|a, \underline{U})$. Then $w^{a, b'}$ is equal to $w_{a, b'}$, an optimal solution to $(\text{Min-Max}|a, b')$. In particular, $w_{a, b'}$ is a GMH contract that implements a , $U(w_{a, b'}, a) = b'$, and $\hat{a}(a, b')$ is an alternate best response to $w_{a, b'}$. Moreover, the Lagrange multipliers $\lambda(a, b')$ and $\delta(a, b')$ in problem $(\text{SAND}|a, \hat{a}(a, b'), b')$ are both positive.*

PROOF. The proof mimics the development in Section 4 of [Ke and Ryan \(2018\)](#) with two key differences. First, [Ke and Ryan \(2018\)](#) do not work with problem $(\text{Min-Max}|a, b')$, but instead with a relaxed problem where $\hat{a}$ is given.⁹ Moreover, the relaxed problem $(P|\hat{a})$ in [Ke and Ryan \(2018\)](#) was defined where the no-jump constraint was an equality. This suffices there because the target action a^* is given. We need more flexibility here and, hence, to follow the logic of [Ke and Ryan \(2018\)](#), we must establish the following claims.

CLAIM 2. *Let $(w_{a, b'}, \hat{a}(a, b'))$ be an optimal solution to $(\text{Min-Max}|a, b')$. Then*

$$U(w_{a, b'}, a) - U(w_{a, b'}, \hat{a}(a, b')) = 0. \quad (38)$$

PROOF. We argue that the Lagrangian multiplier δ^* in [Lemma 3](#) applied to $(\text{SAND}|a, \hat{a}(a, b'), b')$ is strictly greater than zero. Then complementary slackness ([Lemma 3\(ii-b\)](#)) implies that (38) holds.

Suppose $\delta^* = 0$. This implies that w_{a^*} is the first-best contract, denoted $w^{\text{fb}}(b')$. We want to show that a^* is implemented by $w^{\text{fb}}(b')$. This, in turn, implies that the first-best contract is optimal, contradicting [Assumption 3](#). Let $\hat{a}' \in a^{\text{BR}}(w^{\text{fb}}(b'))$ and observe that

$$\begin{aligned} \text{val}(\text{SAND}|a^*, \hat{a}(a^*), b') &= V(w^{\text{fb}}(b'), a^*) \\ &= \inf_{\hat{a} \in \mathbb{A}} \inf_{\lambda, \delta} \max_{w \geq \underline{w}} \mathcal{L}(w, \lambda, \delta | a^*, \hat{a}, b') \\ &\leq \inf_{\lambda, \delta} \max_{w \geq \underline{w}} \mathcal{L}(w, \lambda, \delta | a^*, \hat{a}', b') \\ &= \max_{w \geq \underline{w}} \{V(w, a^*) : U(w, a^*) \geq b', U(w, a^*) - U(w, \hat{a}') \geq 0\} \\ &\leq \max_{w \geq \underline{w}} \{V(w, a^*) : U(w, a^*) \geq b'\} \\ &= V(w^{\text{fb}}(b'), a^*), \end{aligned} \quad (39)$$

where the second equality is by strong duality, the first inequality is by the definition of a minimizer, the third equality is again by strong duality, and the final inequality follows

⁹In that paper ([Ke and Ryan 2018](#)), which determines the optimal choice of $\hat{a}^*$, see the definition of $\hat{a}^*$ in equation (39).

since we have relaxed a constraint. Therefore, all inequalities in the above expression are equalities.

If $U(w^{\text{fb}}(b'), a^*) = U(w^{\text{fb}}(b'), \hat{a}')$, then a^* is a best response to $w^{\text{fb}}(b')$ and we are done. Otherwise, from (39) we must assume $\delta(a^*, \hat{a}') = 0$. This follows by the uniqueness of Lagrangian multipliers (Lemma 3). Therefore, $w^{\text{fb}}(b')$ is the solution to $\arg\max_{w \geq \underline{w}} \{V(w, a^*) : U(w, a^*) \geq b', U(w, a^*) - U(w, \hat{a}') \geq 0\}$ and $U(w^{\text{fb}}(b'), a^*) - U(w^{\text{fb}}(b'), \hat{a}') \geq 0$ is satisfied. Since $\hat{a}' \in a^{\text{BR}}(w^{\text{fb}}(b'))$, we have $a^* \in a^{\text{BR}}(w^{\text{fb}}(b'))$ as desired. $\triangleleft$

The next two claims are adapted from Ke and Ryan (2018). To state them, we need some additional definitions. We let

$$T(x) \equiv \frac{v'(\pi(x) - w^*(x))}{u'(w^*(x))}$$

and

$$R(x) \equiv 1 - \frac{f(x, \hat{a}(a, b'))}{f(x, a)}.$$

Let

$$\mathcal{X}_{\underline{w}}^* = \{x \in \mathcal{X} : w^*(x) = \underline{w}\}.$$

We say two functions φ and ψ with shared domain $\mathcal{X}$ are *co-monotone on the set* $S \subseteq \mathcal{X}$ if φ and ψ are either both nonincreasing or both nondecreasing in S . If φ and ψ are co-monotone on all of $\mathcal{X}$, we simply say that φ and ψ are co-monotone.

CLAIM 3. *If both $T(x)$ and $R(x)$ are co-monotone functions of x on $\mathcal{X} \setminus \mathcal{X}_{\underline{w}}^*$, then w^* is equal to $w_{a, b'}$. Moreover, the Lagrangian multipliers λ and δ associated with the dual of (SAND| $a, \hat{a}(a, b'), b'$) are strictly positive.*

PROOF. This proof is Corollary 5 of Ke and Ryan (2018), setting $\underline{U}$ in that paper to b' . The condition that a is an implementable action and $b' = U(w^{a, \underline{U}}, a)$, where $w^{a, \underline{U}}$ is an optimal solution to (P| $a, \underline{U}$), is required for this proof to hold. $\triangleleft$

The next result is to establish how our assumptions on the output distribution (Assumption 4) guarantee co-monotonicity.

CLAIM 4. *If Assumptions 1–4 hold, then $T(x)$ and $R(x)$ are co-monotone on $\mathcal{X} \setminus \mathcal{X}_{\underline{w}}^*$.*

PROOF. This proof is Lemma 4 of Ke and Ryan (2018). Note that the condition that a be an implementable action and $b' = U(w^{a, \underline{U}}, a)$, where $w^{a, \underline{U}}$ is an optimal solution to (P| $a, \underline{U}$), is required for this proof to hold. Moreover, this also requires Claim 2, where the equality of the no-jump constraint is used to establish (C.14) in the online e-companion of Ke and Ryan (2018). $\triangleleft$

Putting the last two claims together yields [Proposition 6](#). $\square$

An easy implication of the above proposition is that

$$\text{val}(\text{Min-Max}|a, b') = \text{val}(\text{P}|a, b')$$

whenever a is implementable and delivers the agent utility b' in optimality. This will prove to be a useful result in the rest of the proof of [Theorem 1](#). It remains to determine the right implementable a . This is the task of Stage 2.

A.6.2 Analysis of Stage 2 Recall that we are working with a specific $b = U(w^*, a^*)$, where (w^*, a^*) is an optimal solution to (P) (guaranteed to exist by [Assumption 3](#)). The goal of the rest of the proof is to show that $\text{val}(\text{P}) = \text{val}(\text{SAND}|b)$.

We divide this stage of the proof into two further substages. The first substage (Stage 2.1) shows the equivalence between the original problem and a variational max-min-max problem. This intermediate variational problem allows us to leverage the single-dimensional reasoning on display in the proof of [Theorem 1](#) in the single-outcome case in the main body of the paper.

The second substage (Stage 2.2) shows the equivalence between this variational max-min-max and the sandwich problem ([SAND|b](#)).

Stage 2.1. We lighten the notation of Stage 1 and let w_a denote an optimal solution to (Min-Max| a, b) with optimal alternate best response $\hat{a}(a)$ when b is our target agent utility. We construct a variational problem based on w_a as follows. Given $z \in [-1, 1]$, define a set of variations

$$\mathcal{H}(a, z) \equiv \{h \leq \bar{h}(x) : h(x) = 0 \text{ if } w_a(x) = \underline{w} \text{ and } w_a + zh \geq \underline{w} \text{ otherwise}\},$$

where $\bar{h}(x) > w_a(x)$ is sufficiently large, but also where $\int \bar{h}(x)f(x, a) dx < K < \infty$ for some K . We add an additional restriction:

$$\mathcal{M}(a, z) = \left\{ h \in \mathcal{H}(a, z) : \int v'(\pi(x) - w_a(x))h(x)f(x, a) dx \geq 0, \right. \\ \left. \int u'(w_a(x))h(x)f(x, a) dx \geq 0 \right\}.$$

If $h \in \mathcal{M}(a, z)$, then it is not plausible for both the principal and agent to be strictly better off under the variational problem as compared to the original problem. They have a direct conflict of interest in z . This puts us into a situation analogous to the single-outcome case.

We now show the equivalence

$$\text{val}(\text{P}) = \text{val}(\text{Var}|b), \tag{40}$$

where ([Var|b](#)) is the variational optimization problem

$$\max_{a \in \mathbb{A}} \inf_{\hat{a} \in \mathbb{A}} \max_{z \in [-1, 1]} \max_{h \in \mathcal{M}(a, z)} \{V(w_a + zh, a) : (w_a, a) \in \mathcal{W}(\hat{a}, b)\}. \tag{Var|b}$$

The $\leq$ direction of (40) is straightforward since

$$\begin{aligned} & \max_{a \in \mathbb{A}} \inf_{\hat{a} \in \mathbb{A}} \max_{z \in [-1, 1]} \max_{h \in \mathcal{M}(a, z)} \{V(w_a + zh, a) : (w_a, a) \in \mathcal{W}(\hat{a}, b)\} \\ & \geq \inf_{\hat{a} \in \mathbb{A}} \max_{z \in [-1, 1]} \max_{h \in \mathcal{M}(a^*, z)} \{V(w_{a^*} + zh, a^*) : (w_{a^*} + zh, a^*) \in \mathcal{W}(\hat{a}, b)\} \\ & \geq V(w_{a^*}, a^*) = \text{val}(\mathbf{P}), \end{aligned}$$

where the first inequality follows since the optimal action a^* is a feasible choice for a in the outer maximization, the second inequality follows by taking $z = 0$, and the final equality holds from Proposition 6.

It remains to consider the $\geq$ direction of (40). The following claim is analogous Lemma 9 in the proof of Lemma 5.

CLAIM 5. *Given any $\hat{a}$ and a , strong duality holds for (Var|b). That is, for a given $z \in [-1, 1]$,*

$$\begin{aligned} & \max_{h \in \mathcal{M}(a, z)} \{V(w_a + zh, a) : (w_a + zh, a) \in \mathcal{W}(\hat{a}, b)\} \\ & = \inf_{\lambda, \delta, \gamma \geq 0} \max_{h \in \mathcal{H}(a, z)} \mathcal{L}^h(zh, \lambda, \delta, \gamma | a, \hat{a}, b), \end{aligned} \quad (41)$$

where

$$\begin{aligned} \mathcal{L}^h(zh, \lambda, \delta, \gamma | a, \hat{a}, b) &= V(w_a + zh, a) + \lambda[U(w_a + zh, a) - b] \\ &\quad + \delta[U(w_a + zh, a) - U(w_a + zh, \hat{a})] \\ &\quad + \text{sgn}(z)\gamma_1 \int v'(\pi(x) - w_a(x))zh(x)f(x, a)dx \\ &\quad + \text{sgn}(z)\gamma_2 \int u'(w_a(x))zh(x)f(x, a)dx \end{aligned}$$

is the Lagrangian function (which combines the choice of z and h into the product zh since this is how z and h appear in both the objective and the constraints), and $\lambda \geq 0$, $\delta \geq 0$, and $\gamma = (\gamma_1, \gamma_2) \geq 0$ are the Lagrangian multipliers for the remaining constraints that define $\mathcal{M}(a, z)$. Moreover, given that $h^*(\cdot|z)$ solves (41) as a function of z , complementary slackness holds for the optimal choice of $z \in \arg\max_{z \in [-1, 1]} \max_{h \in \mathcal{M}(a, z)} \{V(w_a + zh, a) : (w_a + zh, a) \in \mathcal{W}(\hat{a}, b)\}$; that is,

$$\begin{aligned} & \lambda[U(w_a + zh^*(\cdot|z), a) - b] = 0, \quad \lambda \geq 0, U(w_a + zh^*(\cdot|z), a) - b \geq 0, \\ & \delta[U(w_a + zh, a) - U(w_a + zh^*(\cdot|z), \hat{a})] = 0, \\ & \delta \geq 0, U(w_a + zh^*(\cdot|z), a) \geq U(w_a + zh^*(\cdot|z), \hat{a}), \\ & \gamma_1 \int v'(\pi(x) - w_a(x))h^*(x|z)f(x, a)dx = 0, \\ & \gamma_1 \geq 0, \int v'(\pi(x) - w_a(x))h^*(x|z)f(x, a)dx \geq 0, \\ & \gamma_2 \int u'(w_a(x))h^*(x|z)f(x, a)dx = 0, \quad \gamma_2 \geq 0, \int u'(w_a(x))h^*(x|z)f(x, a)dx \geq 0. \end{aligned}$$

PROOF. By weak duality the $\leq$ direction of (41) is immediate. It remains to consider the $\geq$ direction. For every λ , δ , and γ , $\max_{h \in \mathcal{H}(a, z)} \mathcal{L}^h(zh, \lambda, \delta, \gamma|a, \hat{a}, b)$ is convex in $(\lambda, \delta, \gamma)$. Let $((zh)^*, \lambda^*, \delta^*, \gamma^*)$ denote an optimal solution to the right-hand side of (41). To establish strong duality, we want to show a complementary slackness condition with $(\lambda^*, \delta^*, \gamma^*)$.

The optimization of $\mathcal{L}^h(zh, \lambda, \delta, \gamma|a, \hat{a}, b)$ over zh can be done in a pointwise manner similar to how we approached (SAND| $a, \hat{a}, b$). Given z , by the concavity and monotonicity of v and u , the optimal solution $h(x|z)$ to $\max_{h \in \mathcal{H}(a, \hat{a}, z)} \mathcal{L}^h(zh, \lambda, \delta, \gamma|a, \hat{a}, b)$ must satisfy the following necessary and sufficient conditions:

(i) When $z \geq 0$, $zh(x|z)$ satisfies

$$\left\{ \begin{array}{l} \frac{v'(\pi(x) - w_a(x) - zh(x|z))}{u'(w_a(x) + zh(x|z))} = \left[\lambda + \delta \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) \right] + \frac{\gamma_1 v'(\pi - w_a) + \gamma_2 u'(w_a)}{zu'(w_a(x) + zh(x|z))} \\ \text{if } \frac{v'(\pi(x) - w_a(x))}{u'(w_a(x))} \left(1 - \frac{\gamma_1}{z} \right) \\ \quad < \lambda + \delta \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) + \frac{\gamma_2}{z} \\ \quad \leq \frac{v'(\pi(x) - w_a(x) - z\bar{h}(x))}{u'(w_a(x) + z\bar{h}(x))} - \frac{\gamma_1 v'(\pi - w_a) + \gamma_2 u'(w_a)}{zu'(w_a(x) + z\bar{h}(x))}, \\ h(x|z) = 0 \\ \text{if } \frac{v'(\pi(x) - w_a(x))}{u'(w_a(x))} \left(1 - \frac{\gamma_1}{z} \right) \geq \lambda + \delta \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) + \frac{\gamma_2}{z}, \\ h(x|z) = \bar{h}(x) \\ \text{if } \lambda + \delta \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) + \frac{\gamma_2}{z} > \frac{v'(\pi(x) - w_a(x) - z\bar{h}(x))}{u'(w_a(x) + z\bar{h}(x))} - \frac{\gamma_1 v'(\pi - w_a) + \gamma_2 u'(w_a)}{zu'(w_a(x) + z\bar{h}(x))}. \end{array} \right.$$

(ii) When $z \leq 0$, $zh(x|z)$ satisfies

$$\left\{ \begin{array}{l} \frac{v'(\pi(x) - w_a(x) - zh(x|z))}{u'(w_a(x) + zh(x|z))} = \left[\lambda + \delta \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) \right] + \frac{\gamma_1 v'(\pi - w_a) + \gamma_2 u'(w_a)}{u'(w_a(x) + zh(x|z))} \\ \text{if } \frac{v'(\pi(x) - w_a(x))}{u'(w_a(x))} \left(1 - \frac{\gamma_1}{z} \right) \\ \quad > \lambda + \delta \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) + \frac{\gamma_2}{z} \\ \quad \geq \frac{v'(\pi(x) - w_a(x) - z\bar{h}(x))}{u'(w_a(x) + z\bar{h}(x))} - \frac{\gamma_1 v'(\pi - w_a) + \gamma_2 u'(w_a)}{zu'(w_a(x) + z\bar{h}(x))}, \\ h(x|z) = 0 \\ \text{if } \frac{v'(\pi(x) - w_a(x))}{u'(w_a(x))} \left(1 - \frac{\gamma_1}{z} \right) \leq \lambda + \delta \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) + \frac{\gamma_2}{z}, \\ h(x|z) = \bar{h}(x) \\ \text{if } \lambda + \delta \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) + \frac{\gamma_2}{z} < \frac{v'(\pi(x) - w_a(x) - z\bar{h}(x))}{u'(w_a(x) + z\bar{h}(x))} - \frac{\gamma_1 v'(\pi - w_a) + \gamma_2 u'(w_a)}{zu'(w_a(x) + z\bar{h}(x))}. \end{array} \right.$$

Now we show that, given z , we have the strong duality

$$\max_{h \in \mathcal{H}(a, z)} \{V(w_a + zh, a) : (w_a + zh, a) \in \mathcal{W}(\hat{a}, b)\} = \inf_{\lambda, \delta, \gamma \geq 0} \max_{h \in \mathcal{H}(a, z)} \tilde{\mathcal{L}}^h(zh, \lambda, \delta, \gamma | a, \hat{a}, b),$$

where the Lagrangian is

$$\begin{aligned} \tilde{\mathcal{L}}^h(zh, \lambda, \delta, \gamma | a, \hat{a}, b) &= V(w_a + zh, a) + \lambda[U(w_a + zh, a) - U^*] \\ &\quad + \delta[U(w_a + zh, a) - U(w_a + zh, \hat{a})] \\ &\quad + \gamma_1 \int v'(\pi(x) - w_a(x))h(x)f(x, a) dx \\ &\quad + \gamma_2 \int u'(w_a(x))h(x)f(x, a) dx. \end{aligned}$$

This result follows the uniqueness of $h(x|z)$ as the maximizer of $\tilde{\mathcal{L}}^h(zh, \lambda, \delta, \gamma | a, \hat{a}, b)$ over h . Therefore, the Lagrangian dual function $\psi(\lambda, \delta, \gamma | z) = \max_{h \in \mathcal{H}(a, z)} \tilde{\mathcal{L}}^h(zh, \lambda, \delta, \gamma | a, \hat{a}, b)$ is continuous, differentiable, and convex in $(\lambda, \delta, \gamma)$. This allows us to establish strong duality using similar reasoning as in the proof of [Lemma 3](#).

Let z^* denote the optimal choice of z . We discuss the case where $z^* > 0$. The case where $z^* < 0$ is similar and thus is omitted. In this case, the constraint

$$\int v'(\pi(x) - w_a(x))h(x)f(x, a) dx \geq 0$$

is equivalent to $\int v'(\pi(x) - w_a(x))zh(x)f(x, a) dx \geq 0$ and $\int u'(w_a(x))h(x)f(x, a) dx$ is equivalent to $\int u'(w_a(x))zhf(x, a) dx \geq 0$. Since $h(x|z)$ is uniquely determined, it is continuous in z . Let

$$h^*(x|z^*) \in \arg \max_{h \in \mathcal{H}(a, z^*)} \{V(w_a + z^*h, a) : (w_a + z^*h, a) \in \mathcal{W}(\hat{a}, b)\}$$

be the unique solution to the problem given z^* . Note that $\int v'(\pi(x) - w_a(x))h^*(x|z^*)f(x, a) dx > 0$ and $\int u'(w_a(x))h^*(x|z^*)f(x, a) dx > (1/z^*) \int (u(w_a(x) + z^*h^*(x|z^*)) - u(w_a(x))) \times f(x, a) dx \geq 0$, and

$$\begin{aligned} & - \int v'(\pi(x) - w_a(x) - z^*h^*(x|z^*))h^*(x|z^*)f(x, a) dx \\ & < - \int v'(\pi(x) - w_a(x))h^*(x|z^*)f(x, a) dx \\ & < 0. \end{aligned}$$

Then there must exist Lagrange multipliers $(\lambda^o, \delta^o, \gamma^o)$ such that

$$\begin{aligned} 0 &= \frac{\partial}{\partial z} \mathcal{L}^h(z^*h^*(x|z^*), \lambda^o, \delta^o, \gamma^o | a, \hat{a}, b) \\ &= \int \left(-v'(\pi(x) - w_a(x) - z^*h^*(x|z^*)) \right. \end{aligned}$$

$$\begin{aligned}
& + \left[\lambda^o + \delta^o \left(1 - \frac{f(x, \hat{a})}{f(x, a)} \right) \right] u'(w_a(x) + z^* h^*(x|z^*)) \\
& + \gamma_1^o v'(\pi(x) - w_a(x)) + \gamma_2^o u'(w_a(x)) \Big) h(x|z^*) f(x, a) dx
\end{aligned}$$

and $(\lambda^o, \delta^o, \gamma^o)$ satisfies the complementarity slackness condition. $\square$

The above claim is used to establish another important technical result. The proof is completely analogous to the proof of [Lemma 5](#) in the single-outcome case and thus isomitted.

CLAIM 6. *Let $(a^*, \hat{a}^*, z^*, h^*)$ be an optimal solution to $(\text{Var}|b)$ such that $U(w_{a^*} + z^* h^*, a^*) > b$. Then there exists an $\epsilon > 0$ and optimal solution $(a_\epsilon^*, \hat{a}_\epsilon^*, z_\epsilon^*, h_\epsilon^*)$ such that $U(w_{a_\epsilon^*} + z_\epsilon^* h_\epsilon^*, a_\epsilon^*) = b + \epsilon$ and $V(w_{a^*} + z^* h^*, a^*) = V(w_{a_\epsilon^*} + z_\epsilon^* h_\epsilon^*, a_\epsilon^*)$.*

Via [Claim 6](#), there exists a $b^* \geq b$ and an optimal solution $(\tilde{a}^*, \hat{a}^*, z^*, h^*)$ to $(\text{Var}|b)$ such that $\text{val}(\text{Var}|b) = \text{val}(\text{Var}|b^*)$ and $U(w_{\tilde{a}^*} + z^* h^*, \tilde{a}^*) = b^*$. It then suffices to argue that $\tilde{a}^*$ is implementable (and thus feasible to **(P)**), thus satisfying [\(40\)](#).

To establish implementability, we let $\hat{a}' \in a^{\text{BR}}(w_{\tilde{a}^*} + z^* h^*)$ and claim

$$\begin{aligned}
& V(w_{\tilde{a}^*} + z^* h^*, \tilde{a}^*) \\
& = \max_{z \in [-1, 1]} \max_{h \in \tilde{\mathcal{M}}(\tilde{a}^*, z)} \{ V(w_{\tilde{a}^*} + z^* h^* + zh, \tilde{a}^*) : \\
& \quad (w_{\tilde{a}^*} + z^* h^* + zh, \tilde{a}^*) \in \mathcal{W}(\hat{a}', b^*) \},
\end{aligned} \tag{42}$$

where

$$\tilde{\mathcal{M}}(\tilde{a}^*, z) = \left\{ h \in \tilde{\mathcal{H}}(\tilde{a}^*, z) : \begin{aligned} & \int v'(\pi(x) - w_{\tilde{a}^*}(x) - z^* h^*(x)) h(x) f(x, \tilde{a}^*) dx \geq 0, \\ & \int u'(w_{\tilde{a}^*}(x) + z^* h^*(x)) h(x) f(x, \tilde{a}^*) dx \geq 0 \end{aligned} \right\}$$

and

$$\begin{aligned}
\tilde{\mathcal{H}}(a, z) \equiv \{ & h \leq \bar{h}(x) : h(x) = 0 \text{ if } w_a(x) + z^* h^*(x) + zh(x) = \underline{w} \text{ and} \\
& w_a + z^* h^* + zh \geq \underline{w} \text{ otherwise} \}.
\end{aligned}$$

If [\(42\)](#) holds, then $\tilde{a}^*$ is indeed implementable since $zh = 0$ is a solution to the right-hand side problem. The condition in the right-hand side that $(w_{\tilde{a}^*} + z^* h^*, \tilde{a}^*) \in \mathcal{W}(\hat{a}', b^*)$ implies $U(w_{\tilde{a}^*} + z^* h^*, \tilde{a}^*) \geq U(w_{\tilde{a}^*} + z^* h^*, \hat{a}')$ and so $\tilde{a}^*$ itself must be a best response to $w_{\tilde{a}^*} + z^* h^*$.

To establish [\(42\)](#), note that $\leq$ follows immediately since $(\tilde{a}^*, \hat{a}^*, z^*, h^*)$ solves the left-hand side of [\(40\)](#), whereas in the right-hand side of [\(42\)](#), a particular $\hat{a}$ is chosen (namely

$\hat{a}'$) and there is an additional degree of freedom zh . Next suppose that

$$\begin{aligned} & V(w_{\tilde{a}^*} + z^*h^*, \tilde{a}^*) \\ & < \max_{z \in [-1, 1]} \max_{h \in \tilde{\mathcal{M}}(\tilde{a}^*, z)} \{V(w_{\tilde{a}^*} + z^*h^* + zh, \tilde{a}^*) : \\ & \quad (w_{\tilde{a}^*} + z^*h^* + zh, \tilde{a}^*) \in \mathcal{W}(\hat{a}', b^*)\} \end{aligned} \quad (43)$$

and derive a contradiction.

Let (z^*, h^*) denote an optimal solution to $\max_{z \in [-1, 1]} \max_{h \in \tilde{\mathcal{M}}(\tilde{a}^*, z)} \{V(w_{\tilde{a}^*} + z^*h^* + zh, \tilde{a}^*) : (w_{\tilde{a}^*} + z^*h^* + zh, \tilde{a}^*) \in \mathcal{W}(\hat{a}', b^*)\}$. If (43) holds, then $V(w_{\tilde{a}^*} + z^*h^*, \tilde{a}^*) < V(w_{\tilde{a}^*} + z^*h^* + z^*h^*, \tilde{a}^*)$ and, thus,

$$\begin{aligned} 0 & < V(w_{\tilde{a}^*} + z^*h^* + z^*h^*, \tilde{a}^*) - V(w_{\tilde{a}^*} + z^*h^*, \tilde{a}^*) \\ & \leq -z^{*'} \int h^{*'}(x) v'(\pi(x) - w_{\tilde{a}^*}(x) - z^*h^*(x)) f(x, \tilde{a}^*) dx \end{aligned}$$

since v is concave. Note that $\int h^{*'} v'(\pi - w_{\tilde{a}^*} - z^*h^*) f(x, \tilde{a}^*) dx = 0$ will generate the contradiction $0 < 0$. This further implies $z^{*'} \leq 0$ since $\int h^{*'} v'(\pi - w_{\tilde{a}^*} - z^*h^*) f(x, \tilde{a}^*) dx \geq 0$ by design of the variation set $\tilde{\mathcal{M}}(\tilde{a}^*, z)$. This, in turn, implies $b^* = U(w_{\tilde{a}^*} + z^*h^*, \tilde{a}^*) > U(w_{\tilde{a}^*} + z^*h^* + z^*h^*, \tilde{a}^*)$ since u is concave and $\int h^{*'} u'(w_{\tilde{a}^*} + z^*h^*) f(x, \tilde{a}^*) dx \geq 0$ by design of the variation set $\tilde{\mathcal{M}}(\tilde{a}^*, z)$:

$$\begin{aligned} & U(w_{\tilde{a}^*} + z^*h^* + z^*h^*, \tilde{a}^*) - U(w_{\tilde{a}^*} + z^*h^*, \tilde{a}^*) \\ & = \int [u(w_{\tilde{a}^*}(x) + z^*h^*(x) + z^*h^*(x)) - u(w_{\tilde{a}^*}(x) + z^*h^*(x))] f(x, \tilde{a}^*) dx \\ & < \int z^{*'} h^{*'}(x) u'(w_{\tilde{a}^*}(x) + z^*h^*(x)) f(x, \tilde{a}^*) dx. \\ & \leq 0. \end{aligned}$$

This is a contradiction, since the constraint $(w_{\tilde{a}^*} + z^*h^* + z^*h^*, \tilde{a}^*) \in \mathcal{W}(\hat{a}', b^*)$ implies $U(w_{\tilde{a}^*} + z^*h^* + z^*h^*, \tilde{a}^*) \geq b^*$. This completes Stage 2.1.

Stage 2.2. It remains to show

$$\text{val}(\text{Var}|b) = \text{val}(\text{SAND}|b). \quad (44)$$

Combined with (40), this shows that $\text{val}(\text{P}) = \text{val}(\text{SAND}|b)$, finishing the proof. The direction

$$\begin{aligned} \text{val}(\text{Var}|b) &= \max_{a \in \mathbb{A}} \inf_{\hat{a} \in \mathbb{A}} \max_{z \in [-1, 1]} \max_{h \in \mathcal{M}(a, z)} \{V(w_a + zh, a) : (w_a + zh, a) \in \mathcal{W}(\hat{a}, b)\} \\ &\leq \max_{a \in \mathbb{A}} \inf_{\hat{a} \in \mathbb{A}} \max_{w \geq \underline{w}} \{V(w, a) : (w, a) \in \mathcal{W}(\hat{a}, b)\} = \text{val}(\text{SAND}|b) \end{aligned}$$

follows immediately. It remains to show the $\geq$ direction of (44).

Let $(a^\#, \hat{a}^\#, w_{a^\#})$ be an optimal solution to (SAND| b) that delivers utility $b' \geq b$ to the agent. That is, the constraint $U(w, a) = b'$ is binding in (SAND| b'). We have

$$\text{val}(\text{Var}|b) \geq \inf_{\hat{a} \in \mathbb{A}} \max_{z \in [-1, 1]} \max_{h \in \mathcal{M}(a^\#, z)} \{V(w_{a^\#} + zh, a^\#) : (w_{a^\#} + zh, a^\#) \in \mathcal{W}(\hat{a}, b)\} \quad (45)$$

$$\geq \inf_{\hat{a} \in \mathbb{A}} \max_{z \in [-1, 1]} \max_{h \in \mathcal{M}(a^\#, z)} \{V(w_{a^\#} + zh, a^\#) : (w_{a^\#} + zh, a^\#) \in \mathcal{W}(\hat{a}, b')\} \quad (46)$$

$$= \max_{z \in [-1, 1]} \max_{h \in \mathcal{M}(a^\#, z)} \{V(w_{a^\#} + zh, a^\#) : (w_{a^\#} + zh, a^\#) \in \mathcal{W}(\hat{a}^0, b')\}, \quad (47)$$

where $\hat{a}^0$ is any action in the argmin of the right-hand side of (46). If such an action does not exist, we use a first-order condition following Lemma 4. The details of this case are analogous and thus are omitted. Let $(z^\#, h^\#)$ be in the argmax of the right-hand side of (47). It suffices to show that $\text{val}(\text{SAND}|b)$ is equal to the value of the right-hand side of (47). Observe that $\text{val}(\text{SAND}|b) = \text{val}(\text{SAND}|b')$ and so in the sequel we work with b' without loss.

We argue this in two further substages: (i) We argue that $\text{val}(47) = \text{val}(\text{Min-Max}|a^\#, b^\#)$, where $b^\# = U(w_{a^\#} + z^\#h^\#, a^\#) \geq b'$; for this we use Proposition 6 of Stage 1. (ii) We argue that, in fact $b' = b^\#$. In this case, $\text{val}(\text{Min-Max}|a^\#, b^\#) = \text{val}(\text{Min-Max}|a^\#, b') = \text{val}(\text{SAND}|b')$ since $(a^\#, \hat{a}^\#, w^\#)$ is an optimal solution to (SAND| b'). From (i), this implies $\text{val}(47) = \text{val}(\text{SAND}|b')$. In light of (45)–(47) and the fact $\text{val}(\text{SAND}|b) = \text{val}(\text{SAND}|b')$, this implies $\text{val}(\text{Var}|b) \geq \text{val}(\text{SAND}|b)$ and completes the proof. It remains to establish (i) and (ii) in Stages 2.2.1 and 2.2.2, respectively.

Stage 2.2.1: (i) $\text{val}(47) = \text{val}(\text{Min-Max}|a^\#, b^\#)$. Using similar arguments as in Stage 2.1, we can conclude that $a^\#$ is implemented by $w_{a^\#} + z^\#h^\#$, using the fact that $U(w_{a^\#} + z^\#h^\#, a^\#) = b^\#$ to construct a contradiction.

Given that $w_{a^\#} + z^\#h^\#$ implements $a^\#$ and delivers utility $b^\#$ to the agent, we can apply Proposition 6 to construct an optimal contract $w_{a^\#, b^\#}$ to $(P|a^\#, b^\#)$ with alternate best response $\hat{a}(a^\#, b^\#)$. We then claim that

$$V(w_{a^\#, b^\#}, a^\#) = \text{val}(47). \quad (48)$$

To establish this, we show that $h = w_{a^\#, b^\#} - w_{a^\#}$ belongs to $\mathcal{M}(a^\#, z)$ for $z = 1$. Clearly $w_{a^\#} + h = w_{a^\#, b^\#} \geq \underline{w}$ is satisfied, and $w_{a^\#, b^\#} - w_{a^\#} \leq \bar{h}(x)$ by defining K appropriately large (recall its size was previously left unspecified). Next, we use the concavity of v to see that

$$\begin{aligned} & \int [w_{a^\#, b^\#}(x) - w_{a^\#}(x)] v'(\pi(x) - w_{a^\#}(x)) f(x, a^\#) dx \\ & \geq \int [v(\pi(x) - w_{a^\#}(x)) - v(\pi(x) - w_{a^\#, b^\#}(x))] f(x, a^\#) dx \\ & = \text{val}(\text{SAND}|b) - V(w_{a^\#, b^\#}, a^\#) \\ & \geq \text{val}(\text{SAND}|b) - V(w_{a^\#, b'}, a^\#) \\ & = 0, \end{aligned}$$

since $V(w_{a^\#,b}, a^\#)$ is decreasing in b and using the fact that $b^\# \geq b'$. Next, we note that

$$\begin{aligned} & \int [w_{a^\#,b^\#}(x) - w_{a^\#}(x)] u'(w_{a^\#}(x)) f(x, a^\#) dx \\ & \geq \int [u(w_{a^\#,b^\#}(x)) - u(w_{a^\#}(x))] f(x, a^\#) dx \\ & = b^\# - b' \\ & \geq 0 \end{aligned}$$

by the concavity of u . This shows $h = w_{a^\#,b^\#} - w_{a^\#} \in \mathcal{M}(a^\#, z)$ for $z = 1$. Letting $zh = w_{a^\#,b^\#} - w_{a^\#}$, it is immediate that $w_{a^\#} + zh = w_{a^\#,b^\#} \in \mathcal{W}(\hat{a}^0, b)$. Indeed, $U(w_{a^\#,b^\#}, a^\#) = b^\# \geq b'$ and $U(w_{a^\#,b^\#}, a^\#) - U(w_{a^\#,b^\#}, \hat{a}^0) \geq 0$, since $a^\#$ is implemented by $w_{a^\#,b^\#}$. This implies that $zh = w_{a^\#,b^\#} - w_{a^\#}$ is a feasible choice to (47) and so

$$\text{val}(47) \geq V(w_{a^\#,b^\#}, a^\#).$$

Similarly, since $w_{a^\#} + z^\# h^\#$ is a feasible solution to $(P|a^\#, b^\#)$ (and $w_{a^\#,b^\#}$ is an optimal solution), we get the reverse direction of the above and conclude that

$$V(w_{a^\#} + z^\# h^\#, a^\#) = V(w_{a^\#,b^\#}, a^\#).$$

This yields (48) and completes Stage 2.2.1. This implies that $\hat{a}(a^\#, b^\#)$ can be chosen as $\hat{a}^0$.

Stage 2.2.2: (ii) $b' = b^\#$. It suffices to show that $U(w_{a^\#} + z^\# h^\#, a^\#) = b'$. To do so, we leverage the Lagrangian dual in (41) and argue that the Lagrangian multiplier $\lambda_{z^\#}$ for constraint $U(w_{a^\#} + z^\# h^\#, a^\#) \geq b'$ is strictly positive. Then by complementary slackness, this implies $U(w_{a^\#} + z^\# h^\#, a^\#) = b'$, as required.

Note that $V(w_{a^\#} + z^\# h^\#, a^\#) < V(w_{a^\#}, a^\#)$ (otherwise this already establishes the $\geq$ direction of (40)) and so we have $z^\# > 0$, again using a concavity argument as above. Then $z^\# h^\#$ is uniquely determined by the first-order condition (i) in Claim 5.

Suppose $U(w_{a^\#} + z^\# h^\#, a^\#) > b'$. Then we have $\lambda_{z^\#} = 0$. Then $h^\# = w_{a^\#,b^\#} - w_{a^\#} \neq 0$ implies $\int [w_{a^\#,b^\#}(x) - w_{a^\#}(x)] u'(w_{a^\#}(x)) f(x, a^\#) dx > 0$ and thus $\gamma_2^* = 0$. Moreover, $\text{val}(\text{SAND}|b) > V(w_{a^\#,b^\#}, a^\#)$ implies $\int [w_{a^\#,b^\#}(x) - w_{a^\#}(x)] v'(\pi(x) - w_{a^\#}(x)) f(x, a^\#) dx > 0$, which yields $\gamma_1^* = 0$. Therefore, the first-order condition for $w_{a^\#,b^\#}$ becomes

$$\begin{aligned} \frac{v'(\pi(x) - w_{a^\#,b^\#}(x))}{u'(w_{a^\#,b^\#}(x))} &= \lambda_{z^\#} + \delta_{z^\#} \left(1 - \frac{f(x, \hat{a}^0)}{f(x, a^\#)} \right) \\ &= \delta_{z^\#} \left(1 - \frac{f(x, \hat{a}^0)}{f(x, a^\#)} \right), \quad \text{whenever } w(a^\#, \hat{a}^0, b^\#) > \underline{w}, \end{aligned} \tag{49}$$

where $\lambda_{z^\#}$ and $\delta_{z^\#}$ are the Lagrangian multipliers for the variation problem given $z^\#$. In the case where $\hat{a}_0 \rightarrow a^\#$, Lemma 4 applies and the same structure as (49) holds with the second term equal to $\delta_{z^\#} f_a(x, a^\#)/f(x, a^\#)$. The argument for this case is equivalent

and so we omit it. However, from [Proposition 6](#), we know there is a positive Lagrangian multiplier $\lambda(a^\#, b^\#)$ for optimal contract $w_{a^\#, b^\#}$. By (49) and the fact $w_{a^\#, b^\#}$ is a GMH contract, we have

$$\frac{v'(\pi(x) - w_{a^\#, b^\#}(x))}{u'(w_{a^\#, b^\#}(x))} = \delta_{z^\#} \left(1 - \frac{f(x, \hat{a}^0)}{f(x, a^\#)} \right) = \lambda(a^\#, b^\#) + \delta(a^\#, b^\#) \left(1 - \frac{f(x, \hat{a}^0)}{f(x, a^\#)} \right)$$

for all x such that $w_{a^\#, b^\#}(x) > \underline{w}$. However, if $(1 - f(x, \hat{a}^0)/f(x, a^\#))$ is not a constant for almost all x , the above equalities cannot hold since $\lambda(a^\#, b^\#) > 0$. This contradicts the supposition that $U(w_{a^\#} + z^\# h^\#, a^\#) > b'$ and $\lambda_{z^\#} = 0$.

It only remains to consider the case where $(1 - f(x, \hat{a}^0)/f(x, a^\#))$ is a constant for almost all x such that $w_{a^\#, b^\#}(x) > \underline{w}$. In this case, by the continuity of $v'(\pi(x) - w_{a^\#, b^\#}(x))/u'(w_{a^\#, b^\#}(x))$ in x ($w_{a^\#, b^\#}$ is continuous in x because it is a GMH contract), we know that $v'(\pi(x) - w_{a^\#, b^\#}(x))/u'(w_{a^\#, b^\#}(x))$ is a constant and, thus, characterizes the first-best contract $w(a^\#, b^\#) = w^{\text{fb}}$. Then $w_{a^\#, b^\#}$ implements $a^\#$ and $U(w_{a^\#} + z^\# h^\#, a^\#) = b'$. This completes Stage 2.2.2.

Stage 2.2, Stage 2, and [Theorem 1](#) now follow.

A.7 Proof of [Proposition 1](#)

It suffices to prove the (IR) constraint is binding in (P). Our proof that (IR) is binding is inspired by the proof of Proposition 2 in [Grossman and Hart \(1983\)](#), but adapted to a setting where there are infinitely many (rather than a finite number of) outcomes.

Suppose to the contrary that (w^*, a^*) is an optimal contract where (IR) is not binding; i.e.,

$$U(w^*, a^*) = \underline{U} + \gamma, \quad (50)$$

where $\gamma > 0$. We construct a feasible contract that implements a^* but makes the principal better off, revealing a contradiction.

Under the assumption of the theorem, there exists a $\delta > 0$ such that $w^*(x) > \underline{w} + \delta$ for almost all x . Since u is continuous and increasing, for $\epsilon > 0$ sufficiently small there exists a contract w^ϵ such that

$$w^\epsilon(x) \geq \underline{w} \quad (51)$$

and

$$u(w^\epsilon(x)) = u(w^*(x)) - \epsilon. \quad (52)$$

Observe that for all $a \in \mathbb{A}$,

$$\begin{aligned} U(w^\epsilon, a) &= \int u(w^\epsilon(x)) f(x, a) dx - c(a) \\ &= \int (u(w^*(x)) - \epsilon) f(x, a) dx - c(a) \end{aligned} \quad (53)$$

$$\begin{aligned}
&= \int u(w^*(x))f(x, a) dx - \epsilon \int f(x, a) dx - c(a) \\
&= U(w^*, a) - \epsilon,
\end{aligned}$$

where the first equality is by the definition of U , the second equality is by definition of w^ϵ , the third equality is by the linearity of the integral, and the fourth equality collects terms to form $U(w^*, a)$ and uses the fact $\int f(x, a) dx = 1$ since f is a probability density function.

We are now ready to show that there exists an $\epsilon > 0$ such that (w^ϵ, a^*) is a feasible solution to (P). We already know that w^ϵ satisfies the limited liability constraint for sufficiently small ϵ by (51). We now argue that (IR) and (IC) also hold. For individual rationality, observe that

$$\begin{aligned}
U(w^\epsilon, a^*) &= U(w^*, a^*) - \epsilon \\
&= \underline{U} + \gamma - \epsilon \\
&\geq \underline{U} \quad \text{if } \epsilon < \gamma,
\end{aligned}$$

where the first equality follows from (53) and the second equality uses (50). Since (51) holds for arbitrarily small ϵ , the condition that $\epsilon < \gamma$ can easily be granted.

Finally, for incentive compatibility, observe that for all $a \in \mathbb{A}$,

$$\begin{aligned}
U(w^\epsilon, a^*) - U(w^\epsilon, a) &= [U(w^*, a^*) - \epsilon] - [U(w^*, a) - \epsilon] \\
&= U(w^*, a^*) - U(w^*, a) - \epsilon + \epsilon \\
&\geq 0,
\end{aligned}$$

where the first equality holds from (53) (noting that ϵ is uniform in a). Hence, we conclude that (w^ϵ, a^*) is a feasible solution to (P). Since u is an increasing function, (52) implies $w^\epsilon(x) < w^*(x)$ for all x . Hence, $V(w, a)$ is a decreasing function of w and $w^\epsilon(x) < w^*(x)$, which contradicts the optimality of (w^*, a^*) to (P).

A.8 Proof of Lemma 7

For part (i), since

$$\inf_{\hat{a} \in \mathbb{A}} \inf_{\lambda, \delta \geq 0} \max_{w \geq \underline{w}} \mathcal{L}(w, \lambda, \delta | a, \hat{a}, b) = \inf_{\lambda, \delta \geq 0} \inf_{\hat{a} \in \mathbb{A}} \max_{w \geq \underline{w}} \mathcal{L}(w, \lambda, \delta | a, \hat{a}, b),$$

the desired result follows from the envelope theorem. For part (ii), note that $\inf_{\hat{a}} \mathcal{L}^*(a, \hat{a} | b)$ is continuous and directionally differentiable in a (see, e.g., Corollary 4.4 of Dempe 2002). Since a^* is a maximum, $\frac{\partial}{\partial a^+} (\inf_{\hat{a}} \mathcal{L}^*(a^*, \hat{a} | b)) \leq 0$ and $\frac{\partial}{\partial a^-} (\inf_{\hat{a}} \mathcal{L}^*(a^*, \hat{a} | b)) \geq 0$.

A.9 Proof of Lemma 8

Let b^* be as defined in (19). First, our goal is to show that b^* is tight at optimality, assuming that it exists (we return to existence later in the proof). We first show that

$b^* \leq U(w^*, a^*)$ for all optimal solutions (w^*, a^*) to the original problem (P). Letting $U^* = U(w^*, a^*)$ for some arbitrary optimal solution (w^*, a^*) , we show $b^* \leq U^*$ by arguing that U^* is in the argmin in (19). Our goal is thus to show

$$U^* \in \operatorname{argmin}_{b \geq \underline{U}} \{ \operatorname{val}(\text{SAND}|b) - (P|w(b)) \}. \quad (54)$$

First, observe that

$$\operatorname{val}(P|w(b)) \leq \operatorname{val}(P|b), \quad (55)$$

where $(P|b)$ is defined at the beginning of Section 3. This follows since $(P|w(b))$ considers a problem with a fixed contract $w(b)$ that delivers utility at least b to the agent, whereas $(P|b)$ is an unrestricted version of this problem. Moreover, from Lemma 2 we know that

$$\operatorname{val}(P|b) \leq \operatorname{val}(\text{SAND}|b). \quad (56)$$

Putting (55) and (56) together implies

$$\min_{b \geq \underline{U}} \{ \operatorname{val}(\text{SAND}|b) - \operatorname{val}(P|w(b)) \} \geq 0.$$

With this inequality in hand, we argue that U^* satisfies

$$\operatorname{val}(\text{SAND}|U^*) - \operatorname{val}(P|w(U^*)) = 0, \quad (57)$$

implying our target condition (54). Note that this will also imply that the inner argmin in (19) gives a minimum value of

$$\min_{b \geq \underline{U}} \{ \operatorname{val}(\text{SAND}|b) - \operatorname{val}(P|w(b)) \} = 0. \quad (58)$$

By Theorem 1, we know that (w^*, a^*) is an optimal solution to (P). Also, by Proposition 2, $(w(U^*), a(U^*))$ is an optimal solution to (P). Note, however, that $(w(U^*), a(U^*))$ is also an optimal solution to $(P|w(U^*))$, since feasibility of $(w(U^*), a(U^*))$ to (P) implies $a(U^*) \in a^{\text{BR}}(w(U^*))$. This, in turn, implies $\operatorname{val}(P|w(U^*)) = \operatorname{val}(\text{SAND}|U^*)$ since, as we have just argued, both values are equal to $\operatorname{val}(P)$. This establishes (57) and, hence, we can conclude (54). This shows $b^* \leq U^*$, since b^* is the least element in $\operatorname{argmin}_{b \geq \underline{U}} \{ \operatorname{val}(\text{SAND}|b) - (P|w(b)) \}$. For any tight U^* , since $\operatorname{val}(\text{SAND}|b)$ is a weakly decreasing function of b and $\operatorname{val}(P) = \operatorname{val}(\text{SAND}|U^*)$ for any tight U^* , we have

$$\operatorname{val}(\text{SAND}|b^*) \geq \operatorname{val}(P). \quad (59)$$

Also, by definition (assuming b^* exists), b^* is in the argmin in (19) and so from (58) we know that $\operatorname{val}(P|w(b^*)) = \operatorname{val}(\text{SAND}|b^*)$. However, since $\operatorname{val}(P|w(b^*)) \leq \operatorname{val}(P)$, from (59) we can conclude that $\operatorname{val}(\text{SAND}|b^*) = \operatorname{val}(P)$. In particular, this means that $(w(b^*), a(b^*))$ is an optimal solution to (P). Moreover, from Proposition 6, we know that $U(w(b^*), a(b^*)) = b^*$. Thus, b^* is tight at optimality.

We now show that such a b^* , in fact, exists. Let

$$\hat{b} = \inf\{b \in [\underline{U}, \infty) : \text{val}(\text{SAND}|b) - \text{val}(P|w(b)) = 0\}.$$

For ease of notation let $s(b) = \text{val}(\text{SAND}|b)$ and $t(b) = \text{val}(P|w(b))$. Let B denote the set $\{b \in [\underline{U}, \infty) : s(b) = t(b)\}$ and, thus, $\hat{b}$ is the infimum of B . The goal is to show $\hat{b} \in B$ and, hence, $\hat{b} = b^*$ as defined in (19) using (58). This is achieved by showing that B is closed and bounded below. Clearly B is bounded below by $\underline{U}$. It remains to show closedness. We consider the topological structure of $s(b)$ and $t(b)$. By the theorem of maximum, $s(b)$ is a continuous function of b . Also by the theorem of maximum, $w(b)$ is continuous and $a^{\text{BR}}(b)$ is upper hemicontinuous, and so $t(b)$ is upper semicontinuous. To show that B is closed, consider a sequence b_n in B converging to $\bar{b}$. Since s is a continuous function of b , $\lim_{n \rightarrow \infty} s(b_n) = s(\bar{b})$. Also, since t is upper semicontinuous, we have $\lim_{n \rightarrow \infty} t(b_n) \geq t(\bar{b})$. However, since $t(b) \leq s(b)$ for all b (by (55)), we know that $t(\bar{b}) \leq s(\bar{b})$. Conversely, since $s(b_n) = t(b_n)$, we have $\lim_{n \rightarrow \infty} t(b_n) = \lim_{n \rightarrow \infty} s(b_n) = s(\bar{b})$ and so $s(\bar{b}) \leq t(\bar{b})$. This implies $s(\bar{b}) = t(\bar{b})$, which establishes that B is closed. This completes the proof.

A.10 Proof of Proposition 3

Suppose that for all alternate best responses $\hat{a}$ we have $\hat{a} \geq a$. Observe that when w is a constant function (the same wage for all outputs x), we know that all no-jump constraints

$$U(w, a^*) - U(w, \hat{a}) \geq 0$$

are redundant. Indeed,

$$U(w, a) - U(w, \hat{a}) = c(\hat{a}) - c(a) \geq 0$$

since $\hat{a} \geq a$ and c is an increasing function. Next, observe that when the principal is risk-neutral, the first-best contract is a constant contract. This implies that this constant first-best contract is feasible to (P) and thus optimal, yielding a contradiction.

A.11 Proof of Proposition 4

We now claim that $\text{val}(\text{SAND}|\underline{U}) = \text{val}(\text{FOA})$. First we argue that

$$\text{val}(\text{SAND}|\underline{U}) \geq \text{val}(\text{FOA}). \quad (60)$$

When the first-order approach is valid, we have $\text{val}(\text{FOA}) = \text{val}(P)$. Moreover, by Lemma 2, we also know that $\text{val}(\text{SAND}|\underline{U}) \geq \text{val}(P)$. Putting these together implies (60).

We now turn to showing the reverse inequality of (60); that is,

$$\text{val}(\text{SAND}|\underline{U}) \leq \text{val}(\text{FOA}). \quad (61)$$

By similar reasoning to the proof of Lemma 3, the Lagrangian approach also applies to (FOA), and strong duality holds for (FOA) and its Lagrangian dual (see also Jewitt et al.

2008 for a proof of a setting with certain boundedness assumptions). Let μ^* be the corresponding multiplier for constraint (FOC(a)) in problem (FOA). Let $(a^\#, \hat{a}^\#, w^\#)$ be an optimal solution to (SAND| $\underline{U}$).

If $\mu^* = 0$, then (SAND| $\underline{U}$) has a smaller value than (FOA) by strong duality. This yields (61).

We are left to consider the case where $\mu^* \neq 0$. Suppose $a^\#$ is not a corner solution (similar arguments apply to the corner solution case). If $\mu^* > 0$, we choose some $\hat{a}$ to approach $a^\#$ from below. If $\mu^* < 0$, we choose $\hat{a}$ to approach $a^\#$ from above. Note that the solution $\hat{a}^\#$ is a global minimum (given the choices of the other variables) and so for very small $\epsilon = a^\# - \hat{a}$, for $\hat{a}$ sufficiently close to $a^\#$, we have

$$\begin{aligned}
 \text{val}(\text{SAND}|\underline{U}) &= \inf_{\hat{a}} \inf_{(\lambda, \delta)} \max_{w \geq \underline{w}} \mathcal{L}(w, \lambda, \delta | a^\#, \hat{a}, \underline{U}) \\
 &= \inf_{(\lambda, \delta)} \inf_{\hat{a}} \max_{w \geq \underline{w}} \mathcal{L}(w, \lambda, \delta | a^\#, \hat{a}, \underline{U}) \\
 &\leq \inf_{(\lambda, \delta)} \max_{w \geq \underline{w}} \{V(w, a^\#) + \lambda[U(w, a^\#) - \underline{U}] + \delta \epsilon U_a(w, a^\#) + o(\epsilon)\}.
 \end{aligned} \tag{62}$$

The first equality follows by strong duality of (SAND| $a^\#, \hat{a}, \underline{U}$) with its dual (via Lemma 3). The inequality follows from the mean value theorem. Since $\hat{a}$ approaches $a^\#$ in the direction we chose, we have

$$\begin{aligned}
 &\inf_{(\lambda, \delta)} \max_{w \geq \underline{w}} V(w, a^\#) + \lambda[U(w, a^\#) - \underline{U}] + \delta \epsilon U_a(w, a^\#) \\
 &= \inf_{\lambda} \inf_{\mu \in \mathbb{R}} \max_{w \geq \underline{w}} V(w, a^\#) + \lambda[U(w, a^\#) - \underline{U}] + \mu U_a(w, a^\#) \\
 &\leq \max_{a \in \mathbb{A}} \inf_{\lambda} \inf_{\mu \in \mathbb{R}} \max_{w \geq \underline{w}} V(w, a) + \lambda[U(w, a) - \underline{U}] + \mu U_a(w, a) = \text{val}(\text{FOA}),
 \end{aligned}$$

where we simply redefine $\delta \epsilon = \mu$, without loss of generality. Note that the right-hand side is the statement of the Lagrangian dual of (FOA), and so by strong duality of (FOA) and (62), this implies (61). Combined with (60), this implies $\text{val}(\text{SAND}|\underline{U}) = \text{val}(\text{FOA})$, as required.

REFERENCES

- Araujo, Aloisio and Humberto Moreira (2001), “A general Lagrangian approach for non-concave moral hazard problems.” *Journal of Mathematical Economics*, 35, 17–39. [1426, 1456]
- Bertsekas, Dimitri P. (1999), *Nonlinear Programming*. Athena, Belmont, Massachusetts. [1442]
- Conlon, John R. (2009), “Two new conditions supporting the first-order approach to multisignal principal–agent problems.” *Econometrica*, 77, 249–278. [1425, 1426]

- Dempe, Stephan (2002), *Foundations of Bilevel Programming*. Kluwer Academic Publishers, Dordrecht. [1476]
- Grossman, Sanford J. and Oliver D. Hart (1983), “An analysis of the principal–agent problem.” *Econometrica*, 51, 7–45. [1425, 1426, 1427, 1475]
- Hölmstrom, Bengt (1979), “Moral hazard and observability.” *Bell Journal of Economics*, 10, 74–91. [1440, 1454, 1456]
- Innes, Robert D. (1990), “Limited liability and incentive contracting with ex-ante action choices.” *Journal of Economic Theory*, 52, 45–67. [1425]
- Jewitt, Ian (1988), “Justifying the first-order approach to principal–agent problems.” *Econometrica*, 56, 1177–1190. [1425, 1426, 1452]
- Jewitt, Ian, Ohad Kadan, and Jeroen M. Swinkels (2008), “Moral hazard with bounded payments.” *Journal of Economic Theory*, 143, 59–82. [1443, 1478, 1479]
- Jung, J. Y. and S. K. Kim (2015), “Information space conditions for the first-order approach in agency problems.” *Journal of Economic Theory*, 160, 243–279. [1425]
- Kadan, Ohad, Philip J. Reny, and Jeroen M. Swinkels (2017), “Existence of optimal mechanisms in principal–agent problems.” *Econometrica*, 85, 769–823. [1429, 1453]
- Ke, Rongzhu and Christopher Thomas Ryan (2018), “Monotonicity of optimal contracts without the first-order approach.” *Operations Research*, 66, 1101–1118. [1427, 1459, 1463, 1464, 1465, 1466]
- Kirkegaard, René (2017a), “Moral hazard and the spanning condition without the first-order approach.” *Games and Economic Behavior*, 102, 373–387. [1425]
- Kirkegaard, René (2017b), “A unifying approach to incentive compatibility in moral hazard problems.” *Theoretical Economics*, 12, 25–51. [1425]
- Lasdon, Leon S. (2011), *Optimization Theory for Large Systems*. Dover Publications, Mineola, New York. [1461]
- Mirrlees, James A. (1986), “The theory of optimal taxation.” In *Handbook of Mathematical Economics*, Vol. 3, 1197–1249, Elsevier, Amsterdam. [1426, 1431]
- Mirrlees, James A. (1999), “The theory of moral hazard and unobservable behaviour: Part I.” *The Review of Economic Studies*, 66, 3–21. [1426, 1427, 1430, 1431, 1439, 1441, 1463]
- Rogerson, William P. (1985), “The first-order approach to principal–agent problems.” *Econometrica*, 53, 1357–1367. [1425, 1426, 1427, 1440, 1443, 1452]
- Sinclair-Desgagné, Bernard (1994), “The first-order approach to multi-signal principal–agent problems.” *Econometrica*, 62, 459–465. [1425]

Wang, Wenbin and Shanshan Hu (2016), “Moral hazard with limited liability: The random-variable formulation and optimal contract structures.” Unpublished paper, Indiana University, Department of Operations and Decision Technologies, Kelley School of Business. [[1425](#)]

Co-editor Johannes Hörner handled this manuscript.

Manuscript received 29 April, 2015; final version accepted 19 November, 2017; available online 28 November, 2017.