

Optimal adaptive testing: Informativeness and incentives

RAHUL DEB

Department of Economics, University of Toronto

COLIN STEWART

Department of Economics, University of Toronto

We introduce a learning framework in which a principal seeks to determine the ability of a strategic agent. The principal assigns a test consisting of a finite sequence of tasks. The test is adaptive: each task that is assigned can depend on the agent's past performance. The probability of success on a task is jointly determined by the agent's privately known ability and an unobserved effort level that he chooses to maximize the probability of passing the test. We identify a simple monotonicity condition under which the principal always employs the most (statistically) informative task in the optimal adaptive test. Conversely, whenever the condition is violated, we show that there are cases in which the principal strictly prefers to use less informative tasks.

KEYWORDS. Adaptive testing, dynamic learning, ratcheting, testing experts.

JEL CLASSIFICATION. C44, D82, D83.

1. INTRODUCTION

In this paper, we introduce a learning framework in which a principal seeks to determine the privately known ability of a strategic agent. The principal can assign various tasks to the agent, with different types of the agent varying in their ability to complete a given task.¹ The agent, who is privately informed about his ability, can, through unobservable actions, affect the outcome on each assigned task and thereby influence the information the principal receives. How does strategic behavior by the agent affect the optimal choice of tasks by the principal? When is learning optimized by assigning more informative tasks?

Rahul Deb: rahul.deb@utoronto.ca

Colin Stewart: colinbstewart@gmail.com

This paper was previously circulated with the shorter title “Optimal adaptive testing.” We would like to thank Heski Bar-Isaac, Dirk Bergemann, Rohan Dutta, Amanda Friedenberg, Sean Horan, Johannes Hörner, George Mailath, Marcin Pełski, John Quah, Phil Reny, Larry Samuelson, Ed Schlee, Balázs Szentes, several anonymous referees, and various seminar and conference participants for helpful comments and suggestions. Rya Sciban and Young Wu provided excellent research assistance. We are very grateful for the financial support provided by the Social Sciences and Humanities Research Council of Canada.

¹Learning a worker's ability by observing their performance on differentially informative tasks is a problem that goes back to at least Prescott and Visscher (1980).

Our model can be applied to a number of settings. For example, a manager may choose which tasks to assign to determine whether an employee is suitable for promotion, an employer chooses which questions to ask in a job interview to determine whether a candidate should be hired, and computerized standardized testing assigns questions with the aim of uncovering a student's ability. In each of these scenarios, information is obtained by observing the agent's performance over a sequence of tasks, and the principal's choice of which task to assign may depend on the agent's past performance. Consequently, the agent can, to some extent, control the path of assigned tasks through his performance.

At a more abstract level, our exercise builds on the classic “sequential choice of experiments” problem in statistics (see, for instance, Chapter 14 of [DeGroot 2005](#)). In this problem, a researcher who wants to learn about an unknown parameter has at her disposal a collection of “experiments,” each of which is associated with a different distribution of signals about the parameter. In one formulation, the principal can run a fixed number of experiments, and chooses each experiment sequentially only after observing the outcomes of the preceding ones. A key result in this literature pertains to the case in which one experiment is more informative, in the sense of [Blackwell \(1953\)](#), than all others available to the researcher. In this case, the optimal strategy is independent of the history and simply involves repeatedly drawing from the most informative experiment. We refer to this as Blackwell's result (see Corollary 4.4 in [DeGroot 1962](#)). We introduce strategic behavior into this framework and ask how this strategic behavior by the agent affects the optimal choice of experiments. In particular, does Blackwell's result carry over?

Following the literature on standardized testing, we refer to the optimal task assignment problem that we study as an “adaptive testing” problem. The principal has a fixed number of time periods—corresponding, for instance, to the “tenure clock” in academic institutions or the duration of an interview—over which to evaluate the agent, and a finite collection of different tasks that can be assigned. The agent's probability of success on a particular task depends on his ability (or type) and his choice of action (or effort), neither of which are directly observable to the principal. For instance, the agent may deliberately choose actions that lead to failure if doing so leads to future path of tasks that are more likely to make him look better. Higher actions correspond to a greater probability of success.

The principal first commits to a test. The test begins by assigning the agent a task. Upon seeing the assigned task, the agent chooses his effort level. Depending on the realized success or failure on the first task, the test assigns another task to the agent in the next period, and the agent again chooses his effort. The test continues in this way, with the assigned task in each period possibly depending on the entire history of previous successes and failures. At the end of a fixed number of periods, the test issues a verdict indicating whether the agent passes or fails given the history of tasks and the agent's performance. The principal's goal is to pass the agent if and only if his type belongs to a particular set (which we refer to as the set of good types). As in [Meyer \(1994\)](#), the principal's objective is solely to learn; there are no payoffs associated directly with task completion.

The agent seeks to maximize the probability with which he passes the test. In particular, there are no transfers between the principal and the agent.² Moreover, to focus on the effect of the agent's effort choice on learning, we abstract away from cost-saving incentives by assuming that effort is costless.

A natural benchmark is the optimal test under the assumption that the agent always chooses the highest effort level. Given this strategy, designing the optimal test is essentially a special case of the sequential choice of experiments problem, which can in principle be solved by backward induction (although qualitative properties of the solution are hard to obtain except in the simplest of cases). We refer to this benchmark solution as the *optimal full-effort test* (OFT).

In our strategic environment, Blackwell's result does not hold in general (see Example 2). Our main result (Theorem 2) shows that it does hold if a property we refer to as *group monotonicity* is satisfied, namely, if there does not exist a task at which some bad type has higher ability than some good type. If group monotonicity holds, then it is optimal for the principal always to assign the most informative task and for the agent always to choose the highest effort (in particular, the optimal test coincides with the OFT). Furthermore, the verdict takes a simple form: the agent passes whenever he succeeds at more than some fixed number of tasks.³

We provide a partial converse (Theorem 3) to this result, which indicates that whenever a task violates group monotonicity, there is an environment that includes that task in which always assigning the most informative task is not optimal for the principal. This implies that a task may be the statistically most informative but may not be able to induce maximal learning because assigning it fails to provide the requisite incentives for the strategic agent.

Taken together, these results suggest that in organizations with limited task breadth, where good workers perform better at all tasks (for given levels of effort), managers can optimally learn by assigning the most informative task. However, in organizations that require more task-specific specialization by employees, managers should be concerned about strategic behavior by workers affecting learning.⁴ Similarly, strategic responses must be factored in evaluations of job candidates when they differ in their breadth and level of specialization (such as interviews for academic positions).

In a static setting, the intuition behind our main result is straightforward. Since all types can choose not to succeed on the assigned task, the principal can learn about the agent's type only if success is rewarded with a higher probability of passing the test. In that case, all types choose the highest effort, since doing so maximizes the probability of success. Group monotonicity then ensures that good types have a higher probability of passing than do bad types. Since strategic behavior plays no role, assigning the most informative task is optimal for the principal.

²This assumption is consistent with the observation in Baker et al. (1988) that in many organizations, promotion is the only means used for providing incentives.

³This feature of the optimal test is reminiscent of the optimal contract in the (two outcome version of the) dynamic, pure moral hazard framework of Hölmstrom and Milgrom (1987), where the agent's overall compensation depends only on the number—and not the order—of successes.

⁴Prasad (2009) and Ferreira and Sah (2012) are recent examples of models where workers can be either generalists or specialists.

The dynamic setting is complicated by the fact that the agent must consider how his performance on each task affects the subsequent tasks that will be assigned: he may have an incentive to perform poorly on a task if doing so makes the remainder of the test easier, and thereby increases the ultimate probability of passing. For example, in job interviews, despite it reflecting badly on him, an interviewee may want to deliberately feign ignorance on a particular question, fearing that the line of inquisition that would otherwise follow would be more damaging. [Milgrom and Roberts \(1992\)](#) (see Chapter 7) document strategic shirking in organizations where an employee's own past performance is used as a benchmark for evaluation. In our model, workers are not judged relative to their past performance; however, strategic choice of effort can be used to influence future task assignment and, ultimately, the likelihood of promotion.

It is worth stressing that in our model, even with group monotonicity, there are cases in which some types choose not to succeed on certain tasks in the optimal test (see Example 4). If, however, there is one task q that is more informative than the others, then this turns out not to be an issue. Given any test that, at some histories, assigns tasks other than q , we show that one can recursively replace each of those tasks with q together with a randomized continuation test in a way that does not make the principal worse off. While this procedure resembles Blackwell garbling in the statistical problem, in our case one must be careful to consider how each such change affects the agent's incentives; group monotonicity ensures that any change in the agent's strategy resulting from these modifications to the test can only improve the principal's payoff.

In Section 6, we consider optimal testing when tasks are not comparable in terms of informativeness. We show that under group monotonicity, the OFT is optimal when the agent has only two types (Theorem 4). However, when there are more than two types, this result does not hold: Example 4 shows that even if high effort is always optimal for the agent in the OFT, the principal may be able to do better by inducing some types to shirk. Example 5 and the examples in Appendix B demonstrate a wealth of possibilities (even with group monotonicity).

In Section 7, we consider several extensions of the model. First, we show that our main result continues to hold if the principal can offer the agent a menu of tests from which the agent chooses one (Theorem 5). We also extend the model to allow for the set of available tasks to vary over time. In this case, the optimal test may induce strategic shirking even if it always assigns the most informative task. Finally, we argue that our main result continues to hold if the principal lacks commitment power.

Related literature

Our model and results are related to several distinct strands of the literature. As in the literature on career concerns (beginning with [Holmström 1999](#)), we explore how unobservable effort choices affect learning about an agent's ability. However, our model differs insofar as there are no monetary transfers and the agent has private information about his ability.⁵ In addition, our model incorporates task assignment by the principal. Perhaps the closest related work in this literature is [Dewatripont et al. \(1999\)](#). They

⁵This latter feature of our model more closely resembles the work on screening job applicants using application fees ([Guasch and Weiss 1980](#), [Nalebuff and Scharfstein 1987](#)).

provide conditions under which the market may prefer a less informative monitoring technology (relating the agent's action to performance variables) to a more informative one and vice versa.

More broadly, while more information is always beneficial in a nonstrategic single-agent setting, it can sometimes be detrimental in multi-agent environments. Examples include oligopolies [Mirman et al. \(1994\)](#) and elections [Ashworth et al. \(2017\)](#). While more information is never harmful to the principal in our setting (since she could always choose to ignore it), our focus is on whether less informative tasks can be used to alter the agent's strategy in a way that generates more information.

Our model provides a starting point for studying how managers assign tasks when they benefit from learning about workers' abilities (for instance, to determine their suitability for important projects). Unlike our setting, dynamic contracting is often modeled with pure moral hazard, where the principal chooses bonus payments so as to generate incentives to exert costly effort (see, for instance, [Rogerson 1985](#), [Hölmstrom and Milgrom 1987](#)). However, there are a few recent exceptions that feature both adverse selection and moral hazard. The works of [Gerardi and Maestri \(2012\)](#) and [Halac et al. \(2016\)](#) differ from ours in focus. In these papers, the principal's goal is to learn an unknown state of the world (not the agent's type) and they characterize the optimal transfer schedule for a single task (whereas we study optimal task allocation when promotions are the only means to provide incentives). [Gershkov and Perry \(2012\)](#) also consider a model with transfers, but in their setting, the principal is concerned primarily with matching the complexity of the tasks (which are not assigned by the principal and are instead drawn independently in each period) and the quality of the agent.

The literature on testing forecasters (for surveys, see [Foster and Vohra 2011](#), [Olszewski 2015](#)) shares with our model the aim of designing a test to uncover the type of a strategic agent (an "expert"). In that literature, the expert makes probabilistic forecasts about an unknown stochastic process, and the principal seeks to determine whether the expert knows the true probabilities or is completely ignorant. Our model differs in a number of ways; in particular, the principal assigns tasks and the agent chooses an unobservable action that affects the true probabilities.

Finally, our work is related to the literature on multi-armed bandit problems (an overview can be found in [Bergemann and Välimäki 2006](#)), in which a principal chooses in each period which arm to pull—just as, in our model, she chooses which task to assign—and learns from the resulting outcome. The main trade-off is between maximizing short-term payoffs and the long-term gains from learning. Our model can be thought of as a first step toward understanding bandit problems in which a strategic agent can manipulate the information received by the decision-maker.

2. MODEL

A principal (she) is trying to learn the private type of an agent (he) by observing his performance on a sequence of tasks over T periods.⁶ At each period $t \in \{1, \dots, T\}$, she

⁶Note that T is exogenously fixed. If the principal could choose T , she would always (weakly) prefer it to be as large as possible. Thus, an equivalent alternate interpretation is that the principal has up to T periods to test the agent.

assigns the agent a task q_t from a finite set Q of available tasks. We interpret two identical tasks $q_t = q_{t'}$ assigned at time periods $t \neq t'$ as two different tasks of the same difficulty; the agent being able to succeed on one of the tasks does not imply that he is sure to be able to succeed on the other.

Faced with a task $q_t \in Q$, the agent chooses an effort level $a_t \in [0, 1]$; actions in the interior of the interval may be interpreted as randomization between 0 and 1. All actions have the same cost, which we normalize to 0.⁷ We refer to $a_t = 1$ as full effort and refer to any $a_t < 1$ as shirking. Depending on the agent's ability and effort choice, he may either succeed (s) or fail (f) on a given task. This outcome is observed by both the principal and the agent.

Type space

The agent's ability (which stays constant over time) is captured by his privately known type $\theta_i : Q \rightarrow (0, 1)$, which belongs to a finite set $\Theta = \{\theta_1, \dots, \theta_I\}$.⁸ In period t , the probability of a success on a task q_t when the agent chooses effort a_t is $a_t \theta_i(q_t)$.

The type determines the highest probability of success on each task, obtained when the agent chooses full effort. Zero effort implies sure failure.⁹ Note that, as is common in dynamic moral hazard models, the agent's probability of success on a given task is independent of events that occur before t (such as him having faced the same task before).

Before period 1, the principal announces and *commits to* an (adaptive) test. The test determines which task is assigned in each period depending on the agent's performance so far, and the final verdict given the history at the end of period T .

Histories

At the beginning of period t , h_t denotes a nonterminal public history (or simply a history) up to that point. Such a history lists the tasks faced by the agent and the corresponding successes or failures in periods $1, \dots, t-1$. The set of (nonterminal) histories is denoted by $\mathcal{H} = \bigcup_{t=1, \dots, T} (Q \times \{s, f\})^{t-1}$. We write $\mathcal{H}_{T+1} = (Q \times \{s, f\})^T$ for the set of terminal histories.

Similarly, h_t^A denotes a history for the agent describing his information before choosing an effort level in period t . In addition to the information contained in the history h_t , h_t^A also contains the task he currently faces.¹⁰ Thus the set of all histories for the agent is given by $\mathcal{H}^A = \bigcup_{t=1, \dots, T} (Q \times \{s, f\})^{t-1} \times Q$.

For example, $h_3 = \{(q_1, s), (q_2, f)\}$ is the history at the beginning of period 3 in which the agent succeeded on task q_1 in the first period and failed on task q_2 in the second. The corresponding history $h_3^A = \{(q_1, s), (q_2, f), q_3\}$ also includes the task in period 3.

⁷We make the assumption of identical cost across actions to focus purely on learning, as it ensures that strategic action choices are not muddled by cost-saving incentives.

⁸The restriction that $\theta_i(q) \neq 0$ or 1 simplifies some arguments but is not necessary for any of our results.

⁹The agent's ability to fail for sure is not essential, as none of our results are affected by making the lowest possible effort strictly positive.

¹⁰By not including the agent's actions in h_t^A , we are implicitly excluding the possibility that the agent conditions his effort on his own past choices. Allowing for this would only complicate the notation and make no difference for our results.

Deterministic test

A deterministic test $(\mathcal{T}, \mathcal{V})$ consists of functions $\mathcal{T} : \mathcal{H} \rightarrow \mathcal{Q}$ and $\mathcal{V} : \mathcal{H}_{T+1} \rightarrow \{0, 1\}$. Given a history h_t at the beginning of period t , the task q_t assigned to the agent is $\mathcal{T}(h_t)$. At a terminal history h_{T+1} , the verdict $\mathcal{V}(h_{T+1})$ is the probability with which the agent passes the test.

Test

A (random) test ρ is a distribution over deterministic tests.

As mentioned above, the principal commits to the test in advance. Before period 1, a deterministic test is drawn according to ρ and assigned to the agent. The agent knows ρ but does not observe which deterministic test is realized. He can, however, update as the test proceeds based on the sequence of tasks that have been assigned so far.

Note that even if the agent is facing a deterministic test, since the tasks he faces can depend on his stochastic performance so far in the test, he may not be able to perfectly predict which task he will face in subsequent periods.

Strategies

A strategy for type θ_i is given by a mapping $\sigma_i : \mathcal{H}^A \rightarrow [0, 1]$ from histories for the agent to effort choices; given a history h_t^A in period t , the effort in period t is $a_t = \sigma_i(h_t^A)$. We denote the profile of strategies by $\sigma = (\sigma_1, \dots, \sigma_I)$.

Agent's payoff

Regardless of the agent's type, his goal is to pass the test. Accordingly, faced with a deterministic test $(\mathcal{T}, \mathcal{V})$, the payoff of the agent at any terminal history h_{T+1} is the probability with which he passes, which is given by the verdict $\mathcal{V}(h_{T+1})$. Given a test ρ , we denote by $u_i(h; \rho, \sigma_i)$ the expected payoff of type θ_i when using strategy σ_i conditional on reaching history $h \in \mathcal{H}$.

Principal's beliefs

The principal's prior belief about the agent's type is given by $(\pi_1, \dots, \pi_I)$, with π_i being the probability the principal assigns to type θ_i (thus $\pi_i \geq 0$ and $\sum_{i=1}^I \pi_i = 1$). Similarly, for any $h \in \mathcal{H} \cup \mathcal{H}_{T+1}$, $\pi(h) = (\pi_1(h), \dots, \pi_I(h))$ denotes the principal's belief at history h . We assume that each of these beliefs is consistent with Bayes' rule given the agent's strategy; in particular, at the history $h_1 = \emptyset$, $(\pi_1(h_1), \dots, \pi_I(h_1)) = (\pi_1, \dots, \pi_I)$.

Principal's payoff

The principal partitions the set of types Θ into disjoint subsets of good types $\{\theta_1, \dots, \theta_{i^*}\}$ and bad types $\{\theta_{i^*+1}, \dots, \theta_I\}$, where $i^* \in \{1, \dots, I-1\}$. At any terminal history h_{T+1} , she gets a payoff of 1 if the agent passes and has a good type, -1 if the agent passes and has a bad type, and 0 if the agent fails. Therefore, her expected payoff from a deterministic test $(\mathcal{T}, \mathcal{V})$ is given by $\mathbb{E}_{h_{T+1}}[\sum_{i=1}^{i^*} \pi_i(h_{T+1})\mathcal{V}(h_{T+1}) - \sum_{i=i^*+1}^I \pi_i(h_{T+1})\mathcal{V}(h_{T+1})]$, where

the distribution over terminal histories depends on both the test and the agent's strategy.¹¹

One might expect the principal to receive different payoffs depending on the exact type of the agent, not only whether the type is good or bad. All of our results extend to the more general model in which the principal receives a payoff of γ_i from passing type θ_i , and a payoff normalized to 0 from failing any type. Assuming without loss generality that the types are ordered so that $\gamma_i \geq \gamma_{i+1}$ for each i , the cutoff i^* dividing good and bad types then satisfies $\gamma_i \geq 0$ if $i \leq i^*$ and $\gamma_i \leq 0$ if $i > i^*$. The principal's problem with these more general payoffs and prior π is equivalent to the original problem with prior π' given by $\pi'_i = |\gamma_i| \pi_i / \sum_{j=1}^I |\gamma_j| \pi_j$. Since our results are independent of the prior, this transformation allows us to reduce the problem to the simple binary payoffs for passing the agent described above.

Optimal test

The principal chooses and commits to a test that maximizes her payoff subject to the agent choosing his strategy optimally. Facing a test ρ , we write σ_i^* to denote an optimal strategy for type θ_i , that is, a strategy satisfying

$$\sigma_i^* \in \operatorname{argmax}_{\sigma_i} u_i(h_1; \rho, \sigma_i).$$

Note that this implicitly requires the agent to play optimally at all histories occurring with positive probability given the strategy.

Given her prior, the principal solves

$$\max_{\rho} \mathbb{E}_{h_{T+1}} \left[\mathcal{V}(h_{T+1}) \left(\sum_{i=1}^{i^*} \pi_i(h_{T+1}) - \sum_{i=i^*+1}^I \pi_i(h_{T+1}) \right) \right],$$

where the expectation is taken over terminal histories (the distribution of which depends on the test, ρ , and the strategy $\sigma^* = (\sigma_1^*, \dots, \sigma_I^*)$), and the beliefs are updated from the prior using Bayes' rule (wherever possible). To keep the notation simple, we do not explicitly condition the principal's beliefs π on the agent's strategy.

An equivalent and convenient way to represent the principal's problem is to state it in terms of the agent's payoffs as

$$\max_{\rho} \left[\sum_{i=1}^{i^*} \pi_i v_i(\rho) - \sum_{i=i^*+1}^I \pi_i v_i(\rho) \right], \quad (1)$$

where $v_i(\rho) := u_i(h_1; \rho, \sigma_i^*)$ is the expected payoff type θ_i receives from choosing an optimal strategy in the test ρ . Note in particular that whenever some type of the agent has multiple optimal strategies, the principal is indifferent about which one he employs.

¹¹As in Meyer (1994), we want to focus on the principal's optimal learning problem. This is why we abstract away from payoffs associated with task completion.

3. BENCHMARK: THE OPTIMAL NONSTRATEGIC TEST

Our main goal is to understand how strategic effort choice by the agent affects the principal's ability to learn his type. Thus a natural benchmark is the statistical problem in which the agent is assumed to choose full effort at every history.

Formally, in this benchmark, the principal solves the problem

$$\max_{\mathcal{T}, \mathcal{V}} \mathbb{E}_{h_{T+1}} \left[\mathcal{V}(h_{T+1}) \left(\sum_{i=1}^{i^*} \pi_i(h_{T+1}) - \sum_{i=i^*+1}^I \pi_i(h_{T+1}) \right) \right], \quad (2)$$

where the distribution over terminal histories is determined by the test $(\mathcal{T}, \mathcal{V})$ together with the full-effort strategy

$$\sigma_i^N(h^A) = 1 \quad \text{for all } h^A \in \mathcal{H}^A$$

for every i . We refer to the solution $(\mathcal{T}^N, \mathcal{V}^N)$ to this problem as the optimal full-effort test (OFT). Notice that we have restricted attention to deterministic tests; we argue below that this is without loss.

In principle, it is straightforward to solve for the OFT by backward induction. The principal can first choose the optimal task at all period T histories and beliefs along with the optimal verdicts corresponding to the resulting terminal histories. Observe from (2) that the payoff is linear in the verdicts, so that even if randomization of verdicts is allowed, the optimal choice can always be taken to be either 0 or 1. Moreover, there is no benefit in randomizing tasks: if two tasks yield the same expected payoffs, the principal can choose either one.

Once tasks in period T and verdicts have been determined, it remains to derive the tasks in period $T - 1$ and earlier. At any history h_{T-1} , the choice of task determines the beliefs corresponding to success and failure. In either case, the principal's payoff as a function of those beliefs has already been determined above. Hence the principal simply chooses the task that maximizes her expected payoff. This process can be continued all the way to period 1 to determine the optimal sequence of tasks. At each step, by the same argument as in period T , there is no benefit from randomization. Since the principal may be indifferent between tasks at some history and between verdicts at some terminal history, the OFT need not be unique.

This problem is an instance of the general sequential choice of experiments problem from statistics that we describe in the Introduction. The same backward induction procedure can be applied to (theoretically) solve this more general problem. However, it is typically very difficult to explicitly characterize or to describe qualitative properties of the solution, even in relatively simple special cases that fit within our setting [Bradt and Karlin \(1956\)](#).

4. INFORMATIVENESS

Although the sequential choice of experiments problem is difficult to solve in general, there is a prominent special case that allows for a simple solution: the case in which one task is more Blackwell informative than the others.

Blackwell informativeness

We say that a task q is more Blackwell informative than another task q' if there are numbers $\alpha_s, \alpha_f \in [0, 1]$ such that

$$\begin{bmatrix} \theta_1(q) & 1 - \theta_1(q) \\ \vdots & \vdots \\ \theta_I(q) & 1 - \theta_I(q) \end{bmatrix} \begin{bmatrix} \alpha_s & 1 - \alpha_s \\ \alpha_f & 1 - \alpha_f \end{bmatrix} = \begin{bmatrix} \theta_1(q') & 1 - \theta_1(q') \\ \vdots & \vdots \\ \theta_I(q') & 1 - \theta_I(q') \end{bmatrix}. \quad (3)$$

This is the classic notion of informativeness. Essentially, it says that q is more informative than q' if the latter can be obtained by adding noise to—or garbling—the former. Note that Blackwell informativeness is a partial order; it is possible for two tasks not to be ranked in terms of Blackwell informativeness.

A seminal result due to [Blackwell \(1953\)](#) is that, in any static decision problem, regardless of the decision-maker's preferences, she is always better off with information from a more Blackwell informative experiment than from a less informative one. This result carries over to the sequential setting: if there is one experiment that is more Blackwell informative than every other, then it is optimal for the decision-maker always to use that experiment (see Section 14.17 in [DeGroot 2005](#)). Since the OFT is a special case of this more general problem, if there is a task q that is the most Blackwell informative, then $\mathcal{T}^N(h) = q$ at all $h \in \mathcal{H}$. The following theorem is the formal statement of Blackwell's result applied to our context.

THEOREM 1 ([Blackwell 1953](#)). *Suppose there is a task q that is more Blackwell informative than all other tasks $q' \in Q$. Then there is an OFT in which the task q is assigned at every history.*

In our setting, it is possible to strengthen this result because the principal's payoff takes a special form; Blackwell informativeness is a stronger property than what is needed to guarantee that the OFT features only a single task. We use the term “informativeness” (without the additional “Blackwell” qualifier) to describe the weaker property appropriate for our setting.

Informativeness

Let $\theta_G(q, \pi) = \frac{\sum_{i \leq i^*} \pi_i \theta_i(q)}{\sum_{i \leq i^*} \pi_i}$ be the probability, given belief π , that success is observed on task q conditional on the agent being a good type, under the assumption that the agent chooses full effort. Similarly, let $\theta_B(q, \pi) = \frac{\sum_{i > i^*} \pi_i \theta_i(q)}{\sum_{i > i^*} \pi_i}$ be the corresponding probability of success conditional on the agent being a bad type. We say that a task q is more informative than another task q' if, for all beliefs π , there are numbers $\alpha_s(\pi), \alpha_f(\pi) \in [0, 1]$ such that

$$\begin{bmatrix} \theta_G(q, \pi) & 1 - \theta_G(q, \pi) \\ \theta_B(q, \pi) & 1 - \theta_B(q, \pi) \end{bmatrix} \begin{bmatrix} \alpha_s(\pi) & 1 - \alpha_s(\pi) \\ \alpha_f(\pi) & 1 - \alpha_f(\pi) \end{bmatrix} = \begin{bmatrix} \theta_G(q', \pi) & 1 - \theta_G(q', \pi) \\ \theta_B(q', \pi) & 1 - \theta_B(q', \pi) \end{bmatrix}. \quad (4)$$

To see that Blackwell informativeness is the stronger of these two notions, note that any α_s and α_f that satisfy (3) must also satisfy (4) for every belief π . The following example, consisting of three types and two tasks, shows that the converse need not hold.

EXAMPLE 1. Suppose there are three types ($I = 3$) and two tasks, $Q = \{q, q'\}$. Success probabilities if the agent chooses full effort are

$$\begin{array}{ccc}
 & q & q' \\
 \theta_1 & 0.9 & 0.4 \\
 \theta_2 & 0.8 & 0.2 \\
 \theta_3 & 0.2 & 0.1.
 \end{array} \tag{5}$$

The first column corresponds to the probability $\theta_i(q)$ of success on task q , and the second column corresponds to that on task q' . If $i^* = 2$ (so that types θ_1 and θ_2 are good types), q is more informative than q' . Intuitively, this is because the performance on task q is better at differentiating θ_3 from θ_1 and θ_2 . However, if $i^* = 1$, then q is no longer more informative than q' . This is because performance on task q' is better at differentiating θ_1 from θ_2 . Thus, if the principal's belief assigns high probabilities to θ_1 and θ_2 , she can benefit more from task q' , whereas if her belief assigns high probability to types θ_1 and θ_3 , she can benefit more from q . Since Blackwell informativeness is independent of the cutoff i^* , neither q nor q' is more Blackwell informative than the other. $\diamond$

Although weaker than Blackwell's condition (3), informativeness is still a partial order, and in many cases no element of Q is more informative than all others. However, when there exists a most informative task, our main result shows that Blackwell's result continues to hold for the design of the optimal test in our setting, even when the agent is strategic, provided that a natural monotonicity condition is satisfied. A key difficulty in extending the result is that informativeness is defined independently of the agent's actions and, as the examples in Appendix B demonstrate, in some cases the principal can benefit from strategic behavior by the agent.

5. INFORMATIVENESS AND OPTIMALITY

5.1 The optimal test

The following example shows that strategic behavior by the agent can cause Blackwell's result to fail in our model.

EXAMPLE 2. Suppose there are three types ($I = 3$) and one period ($T = 1$), with $i^* = 2$. There are two tasks, $Q = \{q, q'\}$, with success probabilities given by the matrix

$$\begin{array}{ccc}
 & q & q' \\
 \theta_1 & 0.5 & 0.35 \\
 \theta_2 & 0.2 & 0.5 \\
 \theta_3 & 0.4 & 0.4.
 \end{array}$$

The principal's prior belief is

$$(\pi_1, \pi_2, \pi_3) = (0.3, 0.2, 0.5).$$

Note that task q is more Blackwell informative than q' .¹² If the agent was not strategic, the optimal test would assign task q and verdicts $\mathcal{V}\{(q, s)\} = 0$ and $\mathcal{V}\{(q, f)\} = 1$. In this case, all types would choose $a_1 = 0$, yielding the principal a payoff of 0 (which is the same payoff she would get from choosing either task and $\mathcal{V}\{(q, s)\} = \mathcal{V}\{(q, f)\} = 0$).

Can the principal do better? Assigning task q and reversing the verdicts makes $a_1 = 1$ a best response for all types of the agent but would result in a negative payoff for the principal. Instead, it is optimal for the principal to assign task q' along with verdicts $\mathcal{V}\{(q', s)\} = 1$ and $\mathcal{V}\{(q', f)\} = 0$. Full effort is a best response for all types and this yields a positive payoff. $\diamond$

Notice that in the last example, the types are not ordered in terms of their ability on the tasks the principal can assign. In particular, for each task, the bad type can succeed with higher probability than some good type. This feature turns out to play an important role in determining whether Blackwell's result holds; our main theorem shows that the following condition is sufficient for Blackwell's result to carry over to our model.

Group monotonicity We say that group monotonicity holds if, for every task $q \in Q$, $\theta_i(q) \geq \theta_j(q)$ whenever $i \leq i^* < j$.

This assumption says that the two groups are ordered in terms of ability in a way that is independent of the task: good types are always at least as likely to succeed as bad ones when full effort is chosen.

The proof of our main result builds on a key lemma that, under the assumption of group monotonicity, provides a simple characterization of informativeness that dispenses with the unknowns $\alpha_s(\cdot)$ and $\alpha_f(\cdot)$, and is typically easier to verify than the original definition.

LEMMA 1. *Suppose group monotonicity holds. Then a task q is more informative than q' if and only if*

$$\frac{\theta_i(q)}{\theta_j(q)} \geq \frac{\theta_i(q')}{\theta_j(q')} \quad \text{and} \quad \frac{1 - \theta_j(q)}{1 - \theta_i(q)} \geq \frac{1 - \theta_j(q')}{1 - \theta_i(q')} \quad \text{for all } i \leq i^* \text{ and } j > i^*.$$

Intuitively, a task is more informative if there is a higher relative likelihood that the agent has a good type conditional on a success, and a bad type conditional on a failure. Using this lemma, it is now straightforward to verify that q is more informative than q' in the type space (5) when $i^* = 2$ but not when $i^* = 1$.

We are now in a position to state our main result.

THEOREM 2. *Suppose that there is a task q that is more informative than every other task $q' \in Q$ and that group monotonicity holds. Then any OFT is an optimal test. In particular, it is optimal for the principal to assign task q at all histories and the full-effort strategy σ^N is optimal for the agent.*

¹²The corresponding values of α_s and α_f in (3) are 0.1 and 0.6, respectively.

This result states that the principal cannot enhance learning by inducing strategic shirking through the choice of tasks, a strategy that helps her in Examples 4 and 5. If the principal assigns only the most informative task, it follows from Lemma 2 that she should assign the same verdicts as in the OFT, and the full-effort strategy is optimal for the agent. The optimal verdicts take the form of a cutoff rule where the agent gets a verdict of 1 whenever he succeeds on more than a cutoff number of tasks. Put differently, the only thing that matters is the number of successes and not their order, a feature that can also be found in the dynamic moral hazard environment of [Hölmstrom and Milgrom \(1987\)](#).

While superficially similar, there are critical differences between Theorem 2 and Blackwell's result (Theorem 1). In the latter, where the agent is assumed to always choose the full-effort strategy, the optimality of using the most Blackwell informative task q can be shown constructively by garbling. To see this, suppose that at some history h in the OFT, the principal assigns a task $q' \neq q$, and let α_s and α_f denote the corresponding values that solve (3). In this case, the principal can replace task q' with q and appropriately randomize the continuation tests to achieve the same outcome. More specifically, at the history $\{h, (q, s)\}$, she can choose the continuation test following $\{h, (q', s)\}$ with probability α_s and, with the remaining probability $1 - \alpha_s$, choose the continuation test following $\{h, (q', f)\}$. A similar randomization using α_f can be done at history $\{h, (q, f)\}$.

This construction is not sufficient to yield the result when the agent is strategic. First, if $\alpha_f > 0$, unless all types choose full effort, garbling the continuation test in this way does not generate the same outcome: any type choosing zero effort reaches the continuation test after $\{h, (q', s)\}$ with zero probability before the change and positive probability after the change. Second, replacing the task q' by q and garbling can alter incentives in a way that changes the agent's optimal strategy and, consequently, the principal's payoff. To see this, suppose that full effort is optimal for some type θ_i at $h^A = (h, q')$. This implies that the agent's expected probability of passing the test is higher in the continuation test following $\{h, (q', s)\}$ than in the continuation test following $\{h, (q', f)\}$. Now suppose the principal replaces task q' by q and garbles the continuation tests as described above. Type θ_i may no longer find full effort to be optimal. In particular, if $\alpha_f > \alpha_s$, then zero effort will be optimal after the change, since failure on task q gives a higher likelihood of obtaining the continuation test that he is more likely to pass.¹³ Therefore, the simple garbling argument does not imply Theorem 2. Instead, the proof exploits the structure of informativeness in our particular context captured by Lemma 1, which, when coupled with a backward induction argument, enables us to verify that the continuation tests can be garbled in a way that does not adversely affect incentives.

In our model, i^* is fixed. One could instead ask whether Blackwell's result holds for all i^* when the agent is strategic (that is, for any threshold i^* such that the principal wants to fail the agent if and only if his type has index greater than i^*). It follows from

¹³One can show that if q is more Blackwell informative than q' , group monotonicity implies that the corresponding values of α_f and α_s cannot satisfy $\alpha_f > \alpha_s$. For our weaker notion of informativeness, the relevant values of α_f and α_s may depend on the posterior distribution of types conditional on the history, which in turn depends on the agent's strategy earlier in the test.

Theorem 2 that the result continues to hold if (i) group monotonicity is replaced with the stronger full monotonicity condition that $\theta_i(q) \geq \theta_j(q)$ whenever $i < j$, and (ii) the definition of informativeness is strengthened to require the existence, for *each* partition of the set of types into sets $G = \{\theta_1, \dots, \theta_{i^*}\}$ and $B = \{\theta_{i^*+1}, \dots, \theta_I\}$, of $\alpha_s(\pi)$ and $\alpha_f(\pi)$ satisfying (4).¹⁴

In the nonstrategic benchmark model, Blackwell's result can be strengthened to eliminate less informative tasks even if there is no most informative task. More precisely, if $q, q' \in Q$ are such that q is more informative than q' , then there exists an OFT in which q' is not assigned at any history (and thus any OFT for the set of tasks $Q \setminus \{q'\}$ is also an OFT for the set of tasks Q). The intuition behind this result is essentially the same as for Blackwell's result: whenever a test assigns task q' , replacing it with q and suitably garbling the continuation tests yields the same joint distribution of types and verdicts.

In the strategic setting, this more general result does not hold. For example, there exist cases with one bad type in which zero effort is optimal for the bad type in the first period and full effort is strictly optimal for at least one good type; one such case is described in Example 8 in Appendix B. Letting q denote the task assigned in the first period, adding any task $\tilde{q}$ to the set Q that is easier than q , and assigning $\tilde{q}$ instead of q does not change the optimal action for any type; doing so only increases the payoff of any type that strictly prefers full effort. Since only good types have this preference, such a change increases the principal's payoff. If, in addition, q is more informative than $\tilde{q}$, then the optimal test for the set of tasks $Q \cup \{\tilde{q}\}$ is strictly better for the principal than that for the set Q , which implies that $\tilde{q}$ must be assigned with positive probability at some history, and the generalization of Blackwell's result fails.

5.2 On the structure of the model

While Theorem 2 may seem intuitive, as Example 2 indicates, it does rely on group monotonicity. The following partial converse to Theorem 2 extends the logic of Example 2 to show that, in a sense, group monotonicity is necessary for Blackwell's result to hold in the strategic setting.

THEOREM 3. *Suppose q is such that $\theta_{\bar{i}}(q) < \theta_{\bar{j}}(q)$ for some $\bar{i}$ and $\bar{j}$ such that $\bar{i} \leq i^* < \bar{j}$. Then there exist q' and π such that q is more Blackwell informative than q' , and for each test length T , if $Q = \{q, q'\}$, no optimal test assigns task q at every history $h \in \mathcal{H}$.*

The idea behind this result is that if $\theta_{\bar{i}}(q) < \theta_{\bar{j}}(q)$ and the test always assigns q , type $\bar{j}$ can pass with at least as high a probability as can type $\bar{i}$. When the principal assigns high prior probability to these two types, she is better off assigning a task q' (at least at some histories) for which $\theta_{\bar{i}}(q') > \theta_{\bar{j}}(q')$ (and such a less Blackwell informative q' always exists) so as to advantage the good type.

The next example demonstrates that even if group monotonicity holds, Blackwell's result can also break down if we alter the structure of the agent's payoffs. When all types choose full effort, success on a task increases the principal's belief that the type is good.

¹⁴We are grateful to an anonymous referee for this observation.

Not surprisingly, if some types prefer to fail the test, this can give them an incentive to shirk in a way that overturns Blackwell's result.

EXAMPLE 3. Suppose there are two types ($I = 2$), one good and one bad, and one period ($T = 1$). The principal has two tasks, $Q = \{q, q'\}$, with success probabilities given by the matrix

	q	q'
θ_1	0.9	0.8
θ_2	0.9	0.1.

The principal's prior belief is

$$(\pi_1, \pi_2) = (0.5, 0.5).$$

Compared to the main model, suppose that the principal's payoffs are the same, but the agent's are type-dependent: type θ_1 prefers a verdict of 1 to 0, while type θ_2 has the opposite preference.¹⁵ One interpretation is that verdicts represent promotions to different departments. The principal wants to promote type θ_1 to the position corresponding to verdict 1 and promote type θ_2 to the position corresponding to verdict 0, a preference that the agent shares.

Task q' is trivially more Blackwell informative than task q since the performance on task q (conditional on full effort) conveys no information.¹⁶ Faced with a nonstrategic agent, the optimal test would assign task q' and verdicts $\mathcal{V}\{(q', s)\} = 1$ and $\mathcal{V}\{(q', f)\} = 0$. Faced with a strategic agent, the optimal test is to assign task q and verdicts $\mathcal{V}\{(q, s)\} = 1$ and $\mathcal{V}\{(q, f)\} = 0$. In each of these tests, type θ_1 chooses $a_1 = 1$ and type θ_2 chooses $a_1 = 0$. Thus the probability with which θ_2 gets verdict 0 remains the same, but the probability with which θ_1 gets verdict 1 is higher with the easier task q . $\diamond$

6. NONCOMPARABLE TASKS

In many cases, tasks cannot be ordered by informativeness. What can we say about the design of the optimal test and its relationship to the OFT in general?

The next result shows that when group monotonicity holds, any OFT is an optimal test when there are only two types ($I = 2$); for strategic actions to play an important role, there must be at least three types.

THEOREM 4. *Suppose group monotonicity holds. If $I = 2$, any OFT is an optimal test and makes the full-effort strategy σ^N optimal for the agent.*

To see why the strategy σ^N is optimal for the agent in some optimal test, suppose there is an optimal test in which the good type strictly prefers to shirk at some history h^A . This implies that his expected payoff following a failure on the current task at h^A is higher than that following a success. Now suppose the principal altered the test

¹⁵Given these preferences, it is not clear that full effort is the most relevant benchmark. Our point here is simply to show that, as stated, our result does not carry over to this setting.

¹⁶The corresponding α_s and α_f in (3) are both 0.9.

by replacing the continuation test following a success with that following a failure (including replacing the corresponding verdicts). This would make full effort optimal for both types since the continuation tests no longer depend on success or failure at h^4 . Since the good type chose zero effort before the change, there is no effect on his payoff. Similarly, the bad type's payoff cannot increase: if he strictly preferred full effort before the change, then he is made worse off, and otherwise his payoff is also unchanged. Therefore, this change cannot lower the principal's payoff. A similar argument applies to histories where the bad type prefers to shirk (in which case we can replace the continuation test following a failure with that following a success). Such a construction can be used inductively at all histories where there is shirking.¹⁷

Given this argument, Theorem 4 follows if σ^N is optimal in every OFT. This can be seen using a similar argument to that above, except for the case in which both types strictly prefer to shirk at some history. However, it turns out that this case cannot happen when the continuation tests after both outcomes are chosen optimally.

When there are more than two types, even if group monotonicity holds, there need not be an optimal test in which the fully informative strategy is optimal. The following example shows that, even if the full-effort strategy σ^N is optimal in some OFT, the optimal test may differ; the principal can sometimes benefit from distorting the test relative to the OFT so as to *induce* shirking by some types.

EXAMPLE 4. Suppose there are three types ($I = 3$) and three periods ($T = 3$), with $i^* = 2$ (so that types θ_1 and θ_2 are good types). There are two tasks, $Q = \{q, q'\}$, and the success probabilities are given by the matrix

	q	q'
θ_1	1	0.5
θ_2	0.5	0.5
θ_3	0.5	0.4.

Note that the types are ranked in terms of ability (in particular, group monotonicity holds), and the tasks are ranked in terms of difficulty. The principal's prior belief is

$$(\pi_1, \pi_2, \pi_3) = (0.06, 0.44, 0.5).$$

The OFT $(\mathcal{T}^N, \mathcal{V}^N)$ is represented by the tree in Figure 1. The OFT always assigns the task q' . The agent passes the test if he succeeds at least twice in the three periods. Intuitively, the principal assigns a low prior probability to type θ_1 , and so designs the test to distinguish between types θ_2 and θ_3 , for which q' is better than q . Given that only a single task is used, group monotonicity implies that the optimal verdicts feature a cutoff number of successes required to pass.¹⁸

¹⁷The discussion has ignored the effect of a change following a given period t history on the effort choices at all periods $t' < t$; indeed, earlier actions might change. However, it is straightforward to argue that if a type's payoff goes down at a given history after such a change, the (optimal) payoff is also lower at the beginning of the test.

¹⁸Note that the OFT is not unique in this case since the principal can assign either of the two tasks (keeping the verdicts the same) at histories $\{(q', s), (q', s)\}$ and $\{(q', f), (q', f)\}$.

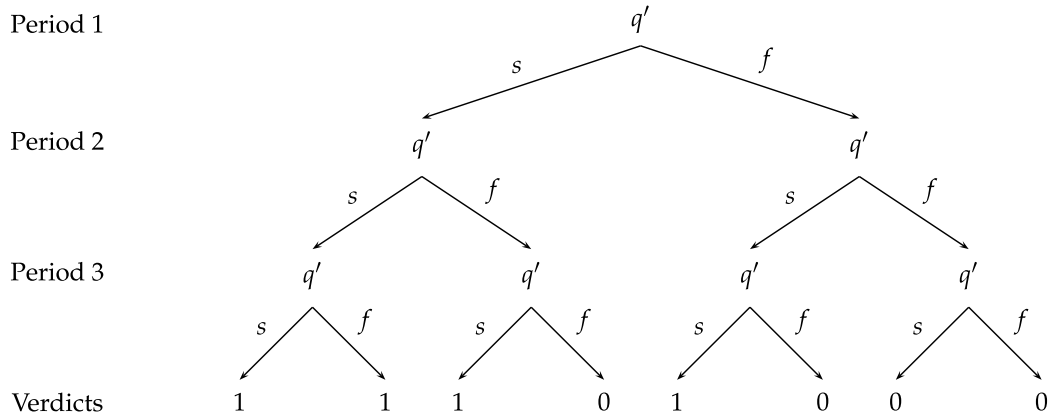


FIGURE 1. An OFT for Example 4. The level of a node corresponds to the time period. Inner nodes indicate the task assigned at the corresponding history, while the leaves indicate the verdicts. For instance, the rightmost node at level 3 corresponds to the period 3 history $h_3 = \{(q', f), (q', f)\}$ and the task assigned by the test at this history is $\mathcal{T}^N(h_3) = q'$. The verdicts following this history are 0 whether he succeeds or fails at this task.

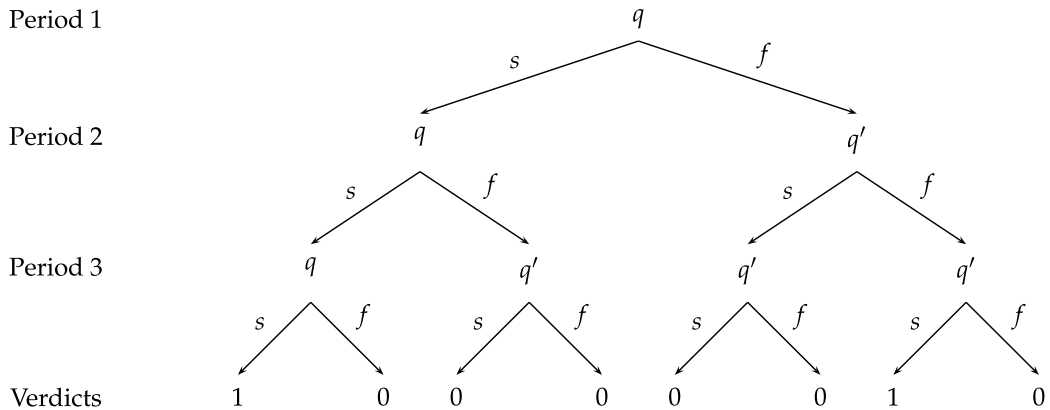


FIGURE 2. An optimal deterministic test for Example 4.

If the principal commits to this test, then the full-effort strategy is optimal for the agent: failure on the task assigned in any period has no effect on the tasks assigned in the future, and merely decreases the probability of passing.

Is this test optimal when the agent is strategic? Consider instead the deterministic test $(\mathcal{T}', \mathcal{V}')$ described by the tree in Figure 2. This alternate test differs from the OFT in several ways. The agent now faces task q instead of q' both in period 1 and at the period 2 history following a success. In addition, the agent can pass only at two of the terminal histories. We argue that this test yields a higher payoff to the principal despite σ^N being an optimal strategy for the agent in test $(\mathcal{T}^N, \mathcal{V}^N)$.

By definition, $(\mathcal{T}', \mathcal{V}')$ can only yield a higher payoff for the principal than does $(\mathcal{T}^N, \mathcal{V}^N)$ if at least one type of the agent chooses to shirk at some history. This is indeed

the case. Since type θ_1 succeeds at task q for sure conditional on choosing full effort, he will choose $a_t = 1$ in each period and pass with probability 1. However, types θ_2 and θ_3 both prefer $a_t = 0$ in periods $t = 1, 2$. Following a success in period 1, two further successes are required at task q to get a passing verdict. In contrast, by choosing the zero effort in the first two periods, the history $\{(q, f), (q', f)\}$ can be reached with probability 1, after which the agent needs only a single success at task q' to pass. Consequently, this shirking strategy yields a higher payoff for types θ_2 and θ_3 .

The difference in payoffs for the three types in $(\mathcal{T}', \mathcal{V}')$ relative to $(\mathcal{T}^N, \mathcal{V}^N)$ are

$$\Delta v_1 = v_1(\mathcal{T}', \mathcal{V}') - v_1(\mathcal{T}^N, \mathcal{V}^N) = 1 - [0.5 * 0.75 + 0.5 * 0.25] = 0.5,$$

$$\Delta v_2 = v_2(\mathcal{T}', \mathcal{V}') - v_2(\mathcal{T}^N, \mathcal{V}^N) = 0.5 - [0.5 * 0.75 + 0.5 * 0.25] = 0,$$

$$\Delta v_3 = v_3(\mathcal{T}', \mathcal{V}') - v_3(\mathcal{T}^N, \mathcal{V}^N) = 0.4 - [0.4 * 0.64 + 0.6 * 0.16] = 0.048.$$

The change in the principal's payoff is

$$\sum_{i=1}^2 \pi_i \Delta v_i - \pi_3 \Delta v_3 = 0.06 * 0.5 - 0.5 * 0.048 > 0,$$

which implies that $(\mathcal{T}^N, \mathcal{V}^N)$ is not the optimal test. In particular, the principal can benefit from the fact that the agent can choose his actions strategically. $\diamond$

The previous example has a flavor of the “ratchet effect.”¹⁹ In the optimal deterministic test, types θ_2 and θ_3 preferred to shirk in the first period, as success on the first task was a signal that the agent might be the highest type θ_1 which, in turn, led to a ratcheting up of the test difficulty (requiring two successes at task q as opposed to one success at q'). The next example shows that such ratcheting might also be a feature of the OFT; specifically, that the full-effort strategy is not always optimal. As the example shows, in response, the principal may be able to improve on the OFT with a different test, even one that induces the same strategy for the agent.

EXAMPLE 5. Suppose there are three types ($I = 3$) and three periods ($T = 3$), with $i^* = 2$. The principal has two different tasks, $Q = \{q, q'\}$, and the success probabilities are

	q	q'
θ_1	1	0.2
θ_2	0.2	0.15
θ_3	0.1	0.01.

The principal's prior belief is

$$(\pi_1, \pi_2, \pi_3) = (0.5, 0.1, 0.4).$$

¹⁹The ratchet effect arises in a variety of different dynamic environments: for instance, in organizations [Milgrom and Roberts \(1992\)](#), in regulatory contexts [Meyer and Vickers \(1997\)](#), and in adverse selection environments [Laffont and Tirole \(1988\)](#).

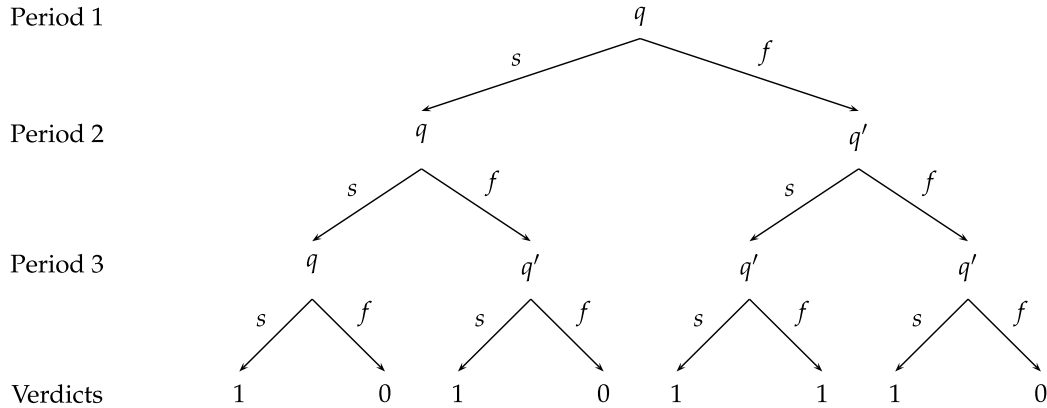


FIGURE 3. An OFT for Example 5.

Figure 3 depicts an OFT $(\mathcal{T}^N, \mathcal{V}^N)$ for this environment. The intuition for the optimality of this test is as follows. The principal has a low prior probability that the agent's type is θ_2 . Task q is effective at distinguishing between types θ_1 and θ_3 as, loosely speaking, their ability difference is larger on that task. If there is a success on q , it greatly increases the belief that the type is θ_1 , and the principal will assign q again. Conversely, if there is a failure on task q (in any period), then the belief assigns zero probability to the agent having type θ_1 . The principal then instead switches to task q' , which is more effective than q at distinguishing between types θ_2 and θ_3 . Since θ_3 has very low ability on q' , a success on this task is a strong signal that the agent's type is not θ_3 , in which case the test issues a pass verdict.

Note that the full-effort strategy σ^N is not optimal for type θ_2 : he prefers to choose action 0 in period 1 and action 1 thereafter. This is because his expected payoff at history $h_2 = \{(q, s)\}$ is $u_2(h_2; \mathcal{T}^N, \mathcal{V}^N, \sigma_2^N) = 0.2 * 0.2 + 0.8 * 0.15 = 0.16$, which is lower than his expected payoff $u_2(h'_2; \mathcal{T}^N, \mathcal{V}^N, \sigma_2^N) = 1 - 0.85 * 0.85 = 0.2775$ at history $h'_2 = \{(q, f)\}$. Therefore, this example demonstrates that the full-effort strategy is not always optimal for the agent in an OFT.²⁰ The ability of the agent to behave strategically benefits the principal since θ_2 is a good type.

An optimal deterministic test $(\mathcal{T}', \mathcal{V}')$ is depicted in Figure 4. Note that this test is identical to $(\mathcal{T}^N, \mathcal{V}^N)$ except that the verdict at terminal history $\{(q, s), (q, f), (q', s)\}$ is 0 as opposed to 1. In this test, types θ_1 and θ_3 choose the full-effort strategy and type θ_2 chooses action 0 in period 1 and action 1 subsequently. Note that the expected payoff of type θ_1 remains unchanged relative to the OFT, but that of type θ_3 is strictly lower. The payoff of type θ_2 is identical to what he receives from optimal play in $(\mathcal{T}^N, \mathcal{V}^N)$. Thus the payoff for the principal from the test $(\mathcal{T}', \mathcal{V}')$ is higher than that from $(\mathcal{T}^N, \mathcal{V}^N)$. $\diamond$

The examples in Appendix B illustrate a range of possibilities for both the optimal test and the OFT. Group monotonicity implies that under the assumption that the agent chooses the full-effort strategy, success on each task raises the principal's belief that the

²⁰Although the OFT is not unique, there is no OFT in this case for which σ^N is optimal.

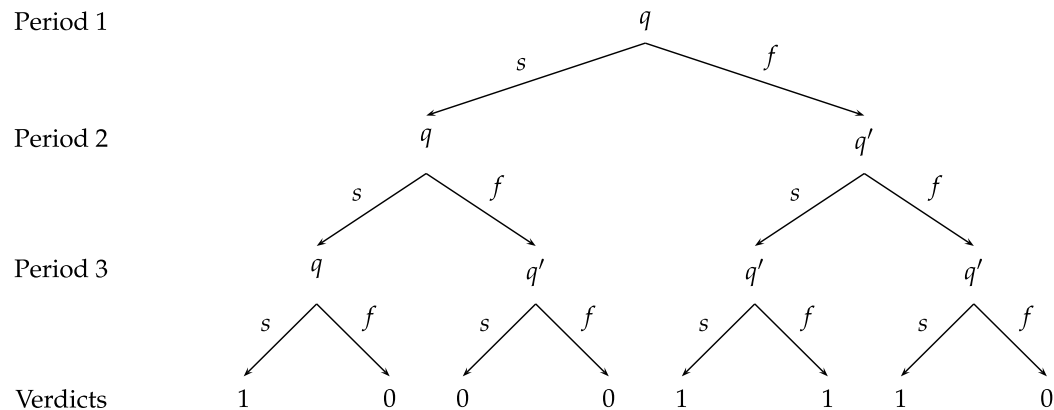


FIGURE 4. An optimal deterministic test for Example 5.

agent's type is good. Nonetheless, because of the adaptive nature of the test, failure on a task can make the remainder of the test easier for some types, as shown by Example 5. Relative to choosing σ^N , strategic behavior by the agent can either help the principal (as in Example 5) or hurt her (as in Example 7). Furthermore, in some cases the full-effort strategy is optimal in the optimal deterministic test but not in the OFT.

Finally, unlike the OFT, for which it suffices to restrict to deterministic tests, there are cases in which there is no deterministic optimal test for the principal when the agent is strategic. Example 8 illustrates one case in which randomizing a verdict strictly benefits the principal and another case in which a test that randomizes tasks is strictly better than any that does not.

7. DISCUSSION

In this section, we consider a number of generalizations of the results in Section 5. Each subsection is separate and does not build on the preceding one.

7.1 Menus of tests

We have so far ignored the possibility that the principal can offer a menu of tests and allow the agent to choose which test to take. While this is not typically observed in the applications we mentioned in the Introduction, it may seem natural from a theoretical perspective. Formally, in this case, the principal offers a menu of M tests $\{\rho_k\}_{k=1}^M$ and each type θ_i of the agent chooses a test ρ_k that maximizes his expected payoff $v_i(\rho_k)$. Although a nontrivial menu could in principle help to screen the different types, our main result still holds.

THEOREM 5. *Suppose there is a task q that is more informative than every other task $q' \in Q$. Then for any OFT, there is an optimal menu consisting only of that test.*

PROOF. In the proof of Theorem 2, we show that any test can be replaced by one where the most informative task q is assigned at all histories and appropriate verdicts can be

chosen so that the payoffs of the good types (weakly) increase and those of the bad types (weakly) decrease. Applying this change to every test in a menu must also increase the good types' payoffs while decreasing those of bad types. Thus we can restrict attention to menus in which every test assigns task q at every history. But then the proof of Lemma 2 shows that replacing any test that is not an OFT with an OFT makes any good type that chooses that test better off and any bad type worse off. Therefore, by the expression for the principal's payoff in (1), replacing every test in the menu with any given OFT cannot make the principal worse off. $\square$

If there is no most informative task, it can happen that offering a nontrivial menu is strictly better for the principal than any single test, as Example 9 in Appendix B shows.

It appears to be very difficult to characterize the optimal menu in general since it involves constructing tests that are themselves complex objects that are challenging to compute. However, without identifying the optimal menu, the following result provides an upper bound on the number of tests that are required: it is always sufficient to restrict to menus containing only as many tests as there are good types. One implication is that nontrivial menus are never beneficial when there is a single good type.

THEOREM 6. *There exists an optimal menu containing at most i^* elements. In particular, if there is a single good type ($i^* = 1$), then there is an optimal menu that consists of a single test.*

PROOF. Suppose the principal offers a menu $\mathcal{M}$, and let $\mathcal{M}'$ denote the subset of $\mathcal{M}$ consisting of the elements chosen by the good types $\theta_1, \dots, \theta_{i^*}$ (so that $\mathcal{M}'$ contains at most i^* elements). If instead of $\mathcal{M}$ the principal offered the menu $\mathcal{M}'$, each good type would continue to choose the same test (or another giving the same payoff), and hence would receive the same payoff as from the menu $\mathcal{M}$. However, the payoff to all bad types must be weakly lower since the set of tests is smaller. Therefore, the menu $\mathcal{M}'$ is at least as good for the principal as $\mathcal{M}$ since it does not affect the probability that any good type passes and weakly decreases the probability that any bad type passes. $\square$

7.2 Time-varying task sets

The main insight from Theorem 2 is that there are simple and intuitive conditions under which the most informative task can generate the required incentives to ensure optimal learning. This result provides the additional insight that the optimal verdict is simple and easy to implement in practice: the agent passes the test whenever he succeeds on sufficiently many tasks. As a consequence, the agent has no incentive to shirk in the optimal test. We now extend the environment to allow for the set of available tasks to vary over time and show that even if the optimal test assigns the most informative task in each period, there may be strategic shirking by the agent.

We now suppose that, in each period t , the principal assigns a task from some nonempty set Q_t that may differ across periods; otherwise, the model is identical to our main model. The next example shows that our main result does not hold in this setting

without additional conditions. One can show, however, that the optimal test always assigns the most informative task if in each period that is also the easiest task for the good types (in the sense that $\theta_i(q_t) \geq \theta_i(q'_t)$ for each $i \leq i^*$, where q_t is the most informative task at time t). The following example also shows that the agent may shirk in the optimal test even when these conditions hold.

EXAMPLE 6. Suppose there are three types ($I = 3$) and three periods ($T = 3$), with $i^* = 2$. In each period t , the principal has only a single task q_t that can be assigned; that is, $Q_t = \{q_t\}$ for each t . The success probabilities are

	q_1	q_2	q_3
θ_1	1	1	0.5
θ_2	1	0.5	1
θ_3	1	0.5	0.5.

The principal's prior belief is

$$(\pi_1, \pi_2, \pi_3) = (0.4, 0.4, 0.2).$$

To describe the OFT $(\mathcal{T}^N, \mathcal{V}^N)$, we only need to define the verdicts as the principal has only one task to assign in each period. Observe that q_1 is completely uninformative (as all types succeed with the same probability) and so does not impact the verdict. It is straightforward to show that the optimal verdict is to pass the agent if and only if he succeeds on at least one of q_2 or q_3 . Since types θ_1 and θ_2 can succeed on at least one of q_2 or q_3 for sure, their expected payoff in the OFT is 1. Type θ_3 passes with probability $1 - 0.5 \times 0.5 = 0.75$. Thus the principal's payoff from the OFT is

$$\sum_{i=1}^2 \pi_i v_i(\mathcal{T}^N, \mathcal{V}^N) - \pi_3 v_3(\mathcal{T}^N, \mathcal{V}^N) = 0.4 + 0.4 - 0.2 \times 0.75 = 0.65.$$

Is this test optimal? Consider instead the verdict function

$$\mathcal{V}(h_{T+1}) = \begin{cases} 1 & \text{if } h_{T+1} = \{(q_1, s), (q_2, s), (q_3, f)\} \\ & \text{or } h_{T+1} = \{(q_1, f), (q_2, f), (q_3, s)\}, \\ 0 & \text{otherwise.} \end{cases} \quad (6)$$

In words, to pass the test, the agent must succeed on exactly one of q_2 or q_3 ; which of these he must succeed on depends on whether he succeeds on q_1 .

How does the agent act in response to these verdicts? Type θ_1 will choose full effort and succeed in periods 1 and 2 for sure, and then shirk and fail in period 3 for sure. Thus, he will pass for sure. Similarly, type θ_2 will shirk and fail in periods 1 and 2 for sure, and then choose full effort and succeed in period 3 for sure. Thus, he too will pass for sure. Finally, type θ_3 is indifferent in period 1, and will choose full effort either in period 2 or period 3, depending on the outcome of q_1 . Thus, he passes with probability 0.5. The

principal's payoff from this test is, therefore,

$$\sum_{i=1}^2 \pi_i v_i(\mathcal{T}^N, \mathcal{V}) - \pi_3 v_3(\mathcal{T}^N, \mathcal{V}) = 0.4 + 0.4 - 0.2 \times 0.5 = 0.7,$$

which is strictly greater than that from the OFT. It follows that the optimal test must induce shirking by some type. This holds even though task assignment is the same as in the OFT and, trivially, the most informative task is assigned in each period. Intuitively, in the second test, the first task is employed as a screening device to allow the two good types to sort into different continuation tests. While it is also possible to screen in this way with time-invariant task sets, doing so is not profitable for the principal.

Finally, we note that if the most informative question is not the easiest one for the good types, it is possible that it will not be used in the optimal test. To see this, suppose we add another task q'_1 to Q_1 so that the period-one task set now becomes $Q'_1 = \{q'_1, q_1\}$. The success probabilities are given by

	q'_1	q_1	q_2	q_3
θ_1	3ε	1	1	0.5
θ_2	2ε	1	0.5	1
θ_3	ε	1	0.5	0.5.

Observe that q'_1 is more Blackwell informative than q_1 .²¹ Suppose the principal assigns task q'_1 instead of q_1 . For small ε , successes almost never occur. Hence, the principal's payoff is close to that obtained from a two-period test using q_2 and q_3 . However, without using q_1 as a screening device, it is no longer possible to achieve the payoff corresponding to the verdict (6). Intuitively, an easier task can be a more effective screening device since it allows good types to sort more fully. $\diamond$

7.3 The role of commitment

Throughout the preceding analysis, we have assumed that the principal can commit in advance to both the history-dependent sequence of tasks and the mapping from terminal histories to verdicts. When the principal cannot commit, her choice of task at each history is determined in equilibrium as a best response to the agent's strategy given the principal's belief. Similarly, the verdicts are chosen optimally at each terminal history depending on the principal's belief (which is also shaped by the agent's strategy). Commitment power benefits the principal (at least weakly) since she can always commit to any equilibrium strategy she employs in the game without commitment (in which case it would be optimal for the agent to choose his equilibrium strategy in response).

If there is a most informative task and group monotonicity holds, then the optimal test can be implemented even without commitment. More precisely, the principal choosing any OFT together with the agent using the strategy σ^N constitutes a sequential equilibrium strategy profile of the game where the principal cannot commit to a test.

²¹Take $\alpha^s = \alpha^f = 1$.

To understand why, note first that the verdicts in this case must correspond directly to the principal's posterior belief at each terminal node, with the agent passing precisely when the principal believes it is more likely that his type is good. Given these verdicts, full effort is optimal in the last period, regardless of what task is assigned, and, hence, by Blackwell's original result, assigning the most informative task is optimal at every history in period T . Given that the same task is assigned at every history in period T , there is no benefit to shirking in period $T - 1$, which implies that assigning the most informative task is again optimal. Working backward in this way yields the result.

In general, optimal tests may not be implementable in the absence of commitment: Example 10 shows how the optimal test may fail to be sequentially rational.

APPENDIX A: PROOFS

We require some additional notation for the proofs. The *length* of a history h_t at the beginning of period t is $|h_t| = t - 1$. We use $S(h_{T+1})$ to denote the number of successes in the terminal history $h_{T+1} \in \mathcal{H}_{T+1}$. Given a history h , the set of all histories of the form $(h, h') \in \mathcal{H}$ is denoted by $\Lambda(h)$ and is referred to as the *subtree* at h . Similarly, we write $\Lambda^A(h)$ for the set of all histories for the agent of the form $(h, h') \in \mathcal{H}^A$. The set of all terminal histories $(h, h') \in \mathcal{H}_{T+1}$ that include h is denoted by $\Gamma(h)$. The length of $\Gamma(h)$ is defined to be $T - |h|$.

For some of the proofs, it is useful to consider tests in which verdicts may be randomized but task assignment is not. A *deterministic test with random verdicts* $(\mathcal{T}, \mathcal{V})$ consists of functions $\mathcal{T} : \mathcal{H} \rightarrow \mathcal{Q}$ and $\mathcal{V} : \mathcal{H}_{T+1} \rightarrow [0, 1]$ (as opposed to the range of $\mathcal{V}$ being $\{0, 1\}$). Note that one can think of any test ρ as randomizing over deterministic tests with random verdicts by combining any tests in the support of ρ that share the same task assignment function $\mathcal{T}$ and defining the randomized verdict function to be the expected verdict conditional on $\mathcal{T}$. In the proofs that follow, we do not distinguish between deterministic tests with or without random verdicts; the meaning is clear from the context.

Given a test ρ and a history (h_t, q_t) for the agent, we write $\text{supp}(h_t, q_t)$ to denote the set of deterministic tests with random verdicts in the support of ρ that generate the history (h_t, q_t) with positive probability if the agent chooses the full-effort strategy.

The following observation is useful for some of the proofs.

OBSERVATION 1. *Given a test ρ , an optimal strategy σ^* for the agent, and a history h , consider an alternate test $\hat{\rho}$ that differs only in the distribution of tasks assigned in the subtree $\Lambda(h)$ and the distribution of verdicts at terminal histories in $\Gamma(h)$. Let $\hat{\sigma}^*$ be an optimal strategy in the test $\hat{\rho}$. Then, for each i , $u_i(h; \hat{\rho}, \hat{\sigma}_i^*) \geq u_i(h; \rho, \sigma_i^*)$ implies $v_i(\hat{\rho}) \geq v_i(\rho)$ and, similarly, $u_i(h; \hat{\rho}, \hat{\sigma}_i^*) \leq u_i(h; \rho, \sigma_i^*)$ implies $v_i(\hat{\rho}) \leq v_i(\rho)$.*

In words, this observation states that if we alter a test at a history h or its subtree $\Lambda(h)$ in a way that the expected payoff of a type increases at h , then the expected payoff also increases at the beginning of the test. This observation is immediate. Consider first the case where $u_i(h; \hat{\rho}, \hat{\sigma}_i^*) \geq u_i(h; \rho, \sigma_i^*)$. Suppose the agent plays the strategy σ_i' such that

$\sigma'_i(h') = \sigma_i^*(h')$ at all histories $h' \notin \Lambda^A(h)$ and $\sigma'_i(h') = \hat{\sigma}_i^*(h')$ at all histories $h' \in \Lambda^A(h)$ on test $\hat{\rho}$. If history h is reached with positive probability, this must yield a weakly higher payoff than playing strategy σ_i^* on test ρ . If history h is reached with probability 0, the payoff of the agent remains the same. Thus the agent can guarantee himself a payoff $u_i(h_1; \hat{\rho}, \sigma'_i) \geq u_i(h_1; \rho, \sigma_i^*)$, which in turn implies that optimal strategy $\hat{\sigma}_i^*$ on $\hat{\rho}$ must yield a payoff at least as high.

The opposite inequality follows from a similar argument. In that case, the agent could only raise his payoff by altering his actions at some histories $h' \notin \Lambda^A(h)$. But if this yielded him a higher payoff, it would contradict the optimality of the strategy σ_i^* .

This observation has a simple implication that we use in what follows: any alteration in a subtree $\Lambda(h)$ that raises the payoffs of good types and lowers the payoffs of bad types leads to a higher payoff for the principal. Formally, if $\hat{\rho}$ differs from ρ only after history h , and if $u_i(h; \hat{\rho}, \hat{\sigma}_i^*) \geq u_i(h; \rho, \sigma_i^*)$ for all $i \leq i^*$ and $u_j(h; \hat{\rho}, \hat{\sigma}_j^*) \leq u_j(h; \rho, \sigma_j^*)$ for all $j > i^*$, then $\hat{\rho}$ yields the principal at least as high a payoff as does ρ (this follows from Observation 1 together with the expression (1) for the principal's payoff).

Proof of Lemma 1

We prove this lemma in two parts. First, we show that q is more informative than q' if and only if, for every π ,

$$\frac{\theta_G(q, \pi)}{\theta_B(q, \pi)} \geq \frac{\theta_G(q', \pi)}{\theta_B(q', \pi)} \quad \text{and} \quad \frac{1 - \theta_B(q, \pi)}{1 - \theta_G(q, \pi)} \geq \frac{1 - \theta_B(q', \pi)}{1 - \theta_G(q', \pi)}. \quad (7)$$

Then we show that the latter condition is equivalent to

$$\frac{\theta_i(q)}{\theta_j(q)} \geq \frac{\theta_i(q')}{\theta_j(q')} \quad \text{and} \quad \frac{1 - \theta_j(q)}{1 - \theta_i(q)} \geq \frac{1 - \theta_j(q')}{1 - \theta_i(q')}$$

for all $i \leq i^*$ and $j > i^*$.

Recall that q is more informative than q' if there is a solution to

$$\begin{bmatrix} \theta_G(q, \pi) & 1 - \theta_G(q, \pi) \\ \theta_B(q, \pi) & 1 - \theta_B(q, \pi) \end{bmatrix} \begin{bmatrix} \alpha_s(\pi) & 1 - \alpha_s(\pi) \\ \alpha_f(\pi) & 1 - \alpha_f(\pi) \end{bmatrix} = \begin{bmatrix} \theta_G(q', \pi) & 1 - \theta_G(q', \pi) \\ \theta_B(q', \pi) & 1 - \theta_B(q', \pi) \end{bmatrix}$$

that satisfies $\alpha_s(\pi), \alpha_f(\pi) \in [0, 1]$. Note that group monotonicity implies that $\theta_G(q, \pi) \geq \theta_B(q, \pi)$. If, for some π , $\theta_G(q, \pi) = \theta_B(q, \pi)$, then such $\alpha_s(\pi)$ and $\alpha_f(\pi)$ exist if and only if $\theta_G(q', \pi) = \theta_B(q', \pi)$. Similarly, (7) holds for the given π if and only if $\theta_G(q', \pi) \leq \theta_B(q', \pi)$. Since $\theta_G(q', \pi) \geq \theta_B(q', \pi)$ by group monotonicity, it follows that $\theta_G(q', \pi) = \theta_B(q', \pi)$. Therefore, when $\theta_G(q, \pi) = \theta_B(q, \pi)$, condition (7) is equivalent to q being more informative than q' .

Now suppose $\theta_G(q, \pi) > \theta_B(q, \pi)$. Solving for $\alpha_s(\pi)$ and $\alpha_f(\pi)$ gives

$$\alpha_s(\pi) = \frac{\theta_G(q', \pi)(1 - \theta_B(q, \pi)) - \theta_B(q', \pi)(1 - \theta_G(q, \pi))}{\theta_G(q, \pi) - \theta_B(q, \pi)},$$

$$\alpha_f(\pi) = \frac{\theta_B(q', \pi)\theta_G(q, \pi) - \theta_G(q', \pi)\theta_B(q, \pi)}{\theta_G(q, \pi) - \theta_B(q, \pi)}.$$

Hence, the condition that $\alpha_s(\pi) \geq 0$ is equivalent to

$$\frac{\theta_G(q', \pi)}{\theta_B(q', \pi)} \geq \frac{1 - \theta_G(q, \pi)}{1 - \theta_B(q, \pi)},$$

which holds because the left-hand side is at least 1 and the right-hand side is less than 1.

The condition that $\alpha_f(\pi) \leq 1$ is equivalent to

$$\frac{\theta_B(q, \pi)}{\theta_G(q, \pi)} \leq \frac{1 - \theta_B(q', \pi)}{1 - \theta_G(q', \pi)},$$

which holds because the left-hand side is less than 1 and the right-hand side is at least 1.

Finally, $\alpha_s(\pi) \leq 1$ is equivalent to

$$\frac{1 - \theta_B(q, \pi)}{1 - \theta_G(q, \pi)} \geq \frac{1 - \theta_B(q', \pi)}{1 - \theta_G(q', \pi)}$$

and $\alpha_f(\pi) \geq 0$ is equivalent to

$$\frac{\theta_G(q, \pi)}{\theta_B(q, \pi)} \geq \frac{\theta_G(q', \pi)}{\theta_B(q', \pi)},$$

which completes the first part of the proof.

We now show the second part. If (7) holds for every π , then given any $i \leq i^*$ and $j > i^*$, taking $\pi_i = \pi_j = \frac{1}{2}$ in (7) gives

$$\frac{\theta_i(q)}{\theta_j(q)} \geq \frac{\theta_i(q')}{\theta_j(q')} \quad \text{and} \quad \frac{1 - \theta_j(q)}{1 - \theta_i(q)} \geq \frac{1 - \theta_j(q')}{1 - \theta_i(q')}.$$

For the converse, observe that

$$\frac{\theta_G(q, \pi)}{\theta_B(q, \pi)} \geq \frac{\theta_G(q', \pi)}{\theta_B(q', \pi)} \iff \sum_{i \leq i^*, j > i^*} \pi_i \pi_j \theta_i(q') \theta_j(q) \left(\frac{\theta_i(q)}{\theta_j(q)} \frac{\theta_j(q')}{\theta_i(q')} - 1 \right) \geq 0,$$

which holds if $\frac{\theta_i(q)}{\theta_j(q)} \geq \frac{\theta_i(q')}{\theta_j(q')}$ whenever $i \leq i^* < j$. Similarly,

$$\begin{aligned} \frac{1 - \theta_B(q, \pi)}{1 - \theta_G(q, \pi)} &\geq \frac{1 - \theta_B(q', \pi)}{1 - \theta_G(q', \pi)} \\ \iff \sum_{i \leq i^*, j > i^*} \pi_i \pi_j (1 - \theta_i(q))(1 - \theta_j(q')) &\left(\frac{1 - \theta_j(q)}{1 - \theta_i(q)} \frac{1 - \theta_i(q')}{1 - \theta_j(q')} - 1 \right) \geq 0, \end{aligned}$$

which holds if $\frac{1 - \theta_j(q)}{1 - \theta_i(q)} \geq \frac{1 - \theta_j(q')}{1 - \theta_i(q')}$ whenever $i \leq i^* < j$. □

Proof of Theorem 2

In what follows, we often refer to a property of verdicts that we term the *cutoff property*. Formally, we say the verdicts in the subtree $\Gamma(h)$ (for a given history h) satisfy the cutoff property (with cutoff k^*) if there exists a number of successes $k^* \in \{0, \dots, T\}$ such that for all terminal histories $h_{T+1} \in \Gamma(h)$, the verdicts satisfy $\mathcal{V}(h_{T+1}) = 1$ whenever $S(h_{T+1}) > k^*$ and satisfy $\mathcal{V}(h_{T+1}) = 0$ whenever $S(h_{T+1}) < k^*$.

The following lemma proves the result for the special case of one task and is useful to prove the general case.

LEMMA 2. *Suppose that $|Q| = 1$ and group monotonicity holds. Then any OFT is an optimal test, and the full-effort strategy σ^N is optimal for the agent.*

PROOF. Since there is only a single task $q \in Q$, a test in this case is simply a deterministic test with random verdicts, which we denote by $(\mathcal{T}, \mathcal{V})$. We begin by stating an observation that is useful for the proof.

OBSERVATION 2. *Suppose that $|Q| = 1$. Suppose in addition that, at some history h , verdicts satisfy the cutoff property with cutoff k^* . Then full effort is optimal for all types at all histories in the subtree $\Lambda(h)$.*

PROOF. This result holds trivially if the length of $\Gamma(h)$ is 1; accordingly, suppose the length is at least 2. First, observe that if this property holds in $\Lambda(h)$, then it also holds in all subtrees of $\Lambda(h)$. Now take any history $h' \in \Lambda(h)$. Consider a terminal history $\{h', (q, f), h''\} \in \Gamma(h)$ following a failure at h' . The verdict at the terminal history $\{h', (q, s), h''\} \in \Gamma(h)$ must be weakly higher, since this is a terminal history with one greater success. Since this is true for all h'' , it implies that any strategy following a failure at h' must yield a weakly lower payoff than if the corresponding strategy was employed after a success. This implies that full effort is optimal at h' . $\square$

We now prove the lemma by induction. The induction hypothesis states that any test $(\mathcal{T}, \mathcal{V})$ of length $T - 1$ that induces shirking (at some history) can be replaced by another test $(\mathcal{T}, \mathcal{V}')$ of the same length in which (i) the full-effort strategy is optimal for every type, (ii) every good type passes with at least as high a probability as in $(\mathcal{T}, \mathcal{V})$, and (iii) every bad type passes with probability no higher than in $(\mathcal{T}, \mathcal{V})$. Therefore, the principal's payoff from test $(\mathcal{T}, \mathcal{V}')$ is at least as high as from $(\mathcal{T}, \mathcal{V})$.

As a base for the induction, consider $T = 1$. If shirking is optimal for some type, it must be that $\mathcal{V}(\{(q, f)\}) \geq \mathcal{V}(\{(q, s)\})$. But then shirking is an optimal action for every type. Changing the verdict function to $\mathcal{V}'(\{(q, s)\}) = \mathcal{V}'(\{(q, f)\}) = \mathcal{V}(\{(q, f)\})$ makes full effort optimal and does not affect the payoff of the agent or the principal.

The induction hypothesis implies that we only need to show that inducing shirking is not strictly optimal for the tester in the first period of a T period test. This induction step is now shown in two separate parts.

Step 1. Consider the two subtrees $\Lambda(\{(q, s)\})$, $\Lambda(\{(q, f)\})$ following success and failure in the first period. For each $\omega \in \{s, f\}$, there exists a number of correct answers

$k_\omega^* \in \{0, \dots, T\}$ such that there are optimal verdicts in the subtree $\Lambda(\{q, \omega\})$ satisfying $\mathcal{V}(h) = 1$ whenever $S(h) > k_\omega^*$ and $\mathcal{V}(h) = 0$ whenever $S(h) < k_\omega^*$ for all $h \in \Gamma(\{q, \omega\})$. Recall that the induction hypothesis states that it is optimal for all types to choose full effort in the subtrees $\Lambda(\{q, s\})$ and $\Lambda(\{q, f\})$. We prove the result for the subtree $\Lambda(\{q, s\})$; an identical argument applies to $\Lambda(\{q, f\})$.

Suppose the statement does not hold. Consider a history $h \in \Lambda(\{q, s\})$ such that the subtree $\Lambda(h)$ is minimal among those in which the statement does not hold (meaning that the statement holds for every proper subtree of $\Lambda(h)$).

Given any optimal verdict function $\mathcal{V}$, let $\underline{k}_s$ and $\bar{k}_s$ denote, respectively, the smallest and largest values of k_s^* for which the statement holds in $\Lambda(\{h, (q, s)\})$. We define $\underline{k}_f$ and $\bar{k}_f$ analogously. If for some optimal $\mathcal{V}$, $\underline{k}_s \leq \bar{k}_f$ and $\underline{k}_f \leq \bar{k}_s$, then there exists k^* for which the statement holds in $\Lambda(\{h, (q, s)\})$ and in $\Lambda(\{h, (q, f)\})$, implying that it holds in $\Lambda(h)$. Therefore, for each optimal $\mathcal{V}$, either $\underline{k}_s > \bar{k}_f$ or $\underline{k}_f > \bar{k}_s$.

Suppose $\underline{k}_s > \bar{k}_f$.²² Let the terminal history $h_{T+1}^s \in \Gamma(\{h, (q, s)\})$ be such that $S(h_{T+1}^s) = \underline{k}_s$ and $\mathcal{V}(h_{T+1}^s) < 1$, and let $h_{T+1}^f \in \Gamma(\{h, (q, f)\})$ be such that $S(h_{T+1}^f) = \bar{k}_f$ and $\mathcal{V}(h_{T+1}^f) > 0$. Note that such terminal histories exist by the minimality and maximality of $\underline{k}_s$ and $\bar{k}_f$, respectively. Let $r = \underline{k}_s - \bar{k}_f$. Let $\tilde{i} \in \arg \min_{i \leq i^*} \theta_i(q)$, and let $\Delta > 0$ be such that

$$\mathcal{V}'(h_{T+1}^s) := \mathcal{V}(h_{T+1}^s) + \Delta \leq 1$$

and

$$\mathcal{V}'(h_{T+1}^f) := \mathcal{V}(h_{T+1}^f) - \frac{\theta_{\tilde{i}}(q)^r}{(1 - \theta_{\tilde{i}}(q))^r} \Delta \geq 0,$$

with one of these holding with equality. Letting $\mathcal{V}'(h) = \mathcal{V}(h)$ for every terminal history $h \notin \{h_{T+1}^s, h_{T+1}^f\}$, changing the verdict function from $\mathcal{V}$ to $\mathcal{V}'$ does not affect the cutoff property in either subtree $\Lambda(\{h, (q, s)\})$ or $\Lambda(\{h, (q, f)\})$. Therefore, by Observation 2, full effort is optimal at all histories in each of these subtrees. In addition, full effort remains optimal at h for all types after the change to $\mathcal{V}'$, because the payoffs of all types have gone up in the subtree $\Lambda(\{h, (q, s)\})$ and down in the subtree $\Lambda(\{h, (q, f)\})$, and, by the induction hypothesis, full effort was optimal at h before the change.

The difference between the expected payoffs for type i at history h (given that the agent follows the full-effort strategy in the subtree $\Lambda(h)$) due to the change in verdicts has the same sign as

$$\Delta \left(\theta_i(q)^r - (1 - \theta_i(q))^r \frac{\theta_{\tilde{i}}(q)^r}{(1 - \theta_{\tilde{i}}(q))^r} \right).$$

By group monotonicity, this last expression is nonnegative if $i \leq \tilde{i}^*$ and nonpositive otherwise. In other words, changing the verdict function to $\mathcal{V}'$ (weakly) raises the payoffs of good types and lowers those of bad types at history h . Therefore, by Observation 1, the principal's payoff does not decrease as a result of this change. Iterating this process at

²²Note that $\underline{k}_s > \bar{k}_f + 1$ implies that it would be optimal for all types to shirk at h , contrary to the induction hypothesis.

other terminal histories $h_{T+1}^s \in \Gamma(\{h, (q, s)\})$ such that $S(h_{T+1}^s) = \underline{k}_s$ and $\mathcal{V}'(h_{T+1}^s) < 1$, and $h_{T+1}^f \in \Gamma(\{h, (q, f)\})$ such that $S(h_{T+1}^f) = \bar{k}_f$ and $\mathcal{V}(h_{T+1}^f) > 0$ eventually leads to an optimal verdict function for which $\underline{k}_s = \bar{k}_f$, as needed.

If $\underline{k}_f > \bar{k}_s$, then all types strictly prefer action 1 at h . To see this, note that for all $\{h, (q, f), h'\} \in \Gamma(\{h, (q, f)\})$, it must be that $\mathcal{V}(\{h, (q, s), h'\}) \geq \mathcal{V}(\{h, (q, f), h'\})$, and this inequality must be strict for all terminal histories where $S(\{h, (q, f), h'\}) = \underline{k}_f - 1$. A similar adjustment to that above can now be done. Let $h_{T+1}^s \in \Gamma(\{h, (q, s)\})$ be such that $S(h_{T+1}^s) = \bar{k}_s$ and $\mathcal{V}(h_{T+1}^s) > 0$, and let $h_{T+1}^f \in \Gamma(\{h, (q, f)\})$ be such that $S(h_{T+1}^f) = \underline{k}_f$ and $\mathcal{V}(h_{T+1}^f) < 1$. Once again, such terminal histories exist by the maximality and minimality of $\bar{k}_s$ and $\underline{k}_f$ respectively. Let $r = \underline{k}_f - \bar{k}_s$, and let $\Delta > 0$ be such that

$$\mathcal{V}'(h_{T+1}^f) = \mathcal{V}(h_{T+1}^f) + \Delta \leq 1$$

and

$$\mathcal{V}'(h_{T+1}^s) = \mathcal{V}(h_{T+1}^s) - \frac{\theta_i(q)^r}{(1 - \theta_i(q))^r} \Delta \geq 0,$$

with one of these holding with equality. As before, this manipulation does not affect the cutoff property at either subtree $\Lambda(\{h, (q, s)\})$ or $\Lambda(\{h, (q, f)\})$ and, therefore, by Observation 2, action 1 is optimal at all histories in each of these subtrees.

Once again, the difference between the expected payoffs for type i at history h (given that the agent follows the full-effort strategy in the subtree $\Lambda(h)$) due to this adjustment has the same sign as

$$\Delta \left(\theta_i(q)^r - (1 - \theta_i(q))^r \frac{\theta_i(q)^r}{(1 - \theta_i(q))^r} \right),$$

which is nonnegative if $i \leq i^*$ and nonpositive if $i > i^*$. Therefore, the principal's payoff does not decrease from these changes, and iterating leads to optimal verdicts satisfying $\underline{k}_f = \bar{k}_s$.

Step 2. Suppose the verdicts at terminal histories $\Gamma(\{(q, s)\})$ and $\Gamma(\{(q, f)\})$ satisfy the above cutoff property, with cutoffs k_s^* and k_f^* , respectively. Then if one type has an incentive to shirk in the first period, so do all other types. Consequently, if all types choose $a_1 = 1$, the proposition follows, or if all types want to shirk, the proposition follows by replacing the test after $\{(q, s)\}$ with the test after $\{(q, f)\}$. This step is straightforward and can be shown by examining the three possible cases. Suppose $k_s^* \leq k_f^*$. Then the verdict at every terminal history $\{(q, s), h\} \in \Gamma(\{(q, s)\})$ is weakly higher than $\{(q, f), h\} \in \Gamma(\{(q, f)\})$ and, hence, $a_1 = 1$ must be optimal for all types. When $k_s^* > k_f^* + 1$, $a_1 = 0$ is optimal for all types. Finally, when $k_s^* = k_f^* + 1$, type i wants to shirk if and only if the sum of the verdicts at terminal histories $\{(q, f), h\} \in \Gamma(\{(q, f)\})$ with $S(\{(q, f), h\}) = k_f^*$ is higher than the sum of the verdicts at terminal histories $\{(q, s), h\} \in \Gamma(\{(q, s)\})$ with $S(\{(q, s), h\}) = k_s^*$ (since each such history occurs with equal probability). This comparison does not depend on i . $\square$

PROOF OF THEOREM 2. In this proof, we proceed backward from period T , altering each deterministic test with random verdicts in the support of ρ in a way that only task q

is assigned without reducing the payoff of the principal. The result then follows from Lemma 2.

Consider first a period T history h_T together with an assigned task q_T . Let

$$v^\omega := \mathbb{E}[\mathcal{V}(\{h_T, (q_T, \omega)\}) | (\mathcal{T}, \mathcal{V}) \in \text{supp}(h_T, q_T)]$$

be the expected verdict following the outcome $\omega \in \{s, f\}$ taken with the respect to the set of possible deterministic tests with random verdicts that the agent could be facing.

Suppose first that $v^f \geq v^s$. Then every type finds shirking optimal at (h_T, q_T) and gets expected verdict v^f . Replacing each deterministic test with random verdicts $(\mathcal{T}, \mathcal{V}) \in \text{supp}(h_T, q_T)$ with another $(\mathcal{T}', \mathcal{V}')$ that is identical except that $\mathcal{T}'(h_T) = q$ and $\mathcal{V}'(\{h_T, (q, s)\}) = \mathcal{V}'(\{h_T, (q, f)\}) = v^f$ does not alter the principal's or the agent's payoff, and makes action 1 optimal at h_T .

Now suppose that $v^s > v^f$, so that action 1 is optimal for all types of the agent. Let $\beta_1 := \max_{i' \leq i^*} \frac{\theta_{i'}(q_T)}{\theta_{i'}(q)}$. If $\beta_1 \leq 1$, we replace each $(\mathcal{T}, \mathcal{V}) \in \text{supp}(h_T, q_T)$ with $(\mathcal{T}', \mathcal{V}')$ that is identical except that $\mathcal{T}'(h_T) = q$, $\mathcal{V}'(\{h_T, (q, s)\}) = \beta_1 v^s + (1 - \beta_1) v^f$, and $\mathcal{V}'(\{h_T, (q, f)\}) = v^f$. The change in expected payoff at history h_T is given by

$$\begin{aligned} & \theta_i(q)(\beta_1 v^s + (1 - \beta_1) v^f) + (1 - \theta_i(q)) v^f - (\theta_i(q_T) v^s + (1 - \theta_i(q_T)) v^f) \\ &= \theta_i(q_T)(v^s - v^f) \left(\frac{\theta_i(q)}{\theta_i(q_T)} \beta_1 - 1 \right) \\ &= \theta_i(q_T)(v^s - v^f) \left(\frac{\theta_i(q)}{\theta_i(q_T)} \max_{i' \leq i^*} \left\{ \frac{\theta_{i'}(q_T)}{\theta_{i'}(q)} \right\} - 1 \right). \end{aligned}$$

Since $v^s - v^f > 0$, it follows from Lemma 1 that the above result is nonnegative for $i \leq i^*$ and nonpositive for $i > i^*$.

Now suppose $\beta_1 > 1$. Let $\beta_2 := 1 - \max_{i' \leq i^*} \frac{\theta_{i'}(q_T) - \theta_{i'}(q)}{1 - \theta_{i'}(q)}$ and observe that $0 \leq \beta_2 \leq 1$ (with the latter inequality following from the assumption that $\beta_1 > 1$). In this case, we replace each $(\mathcal{T}, \mathcal{V}) \in \text{supp}(h_T, q_T)$ with $(\mathcal{T}', \mathcal{V}')$ that is identical except that $\mathcal{T}'(h_T) = q$, $\mathcal{V}'(\{h_T, (q, s)\}) = v^s$ and $\mathcal{V}'(\{h_T, (q, f)\}) = \beta_2 v^f + (1 - \beta_2) v^s$. The change in expected payoff at history h_T is given by

$$\begin{aligned} & \theta_i(q) v^s + (1 - \theta_i(q)) (\beta_2 v^f + (1 - \beta_2) v^s) - (\theta_i(q_T) v^s + (1 - \theta_i(q_T)) v^f) \\ &= (\theta_i(q_T) - \theta_i(q)) (v^s - v^f) \left(\frac{1 - \theta_i(q)}{\theta_i(q_T) - \theta_i(q)} \max_{i' \leq i^*} \frac{\theta_{i'}(q_T) - \theta_{i'}(q)}{1 - \theta_{i'}(q)} - 1 \right). \end{aligned} \quad (8)$$

Note that for any i and i' , $\frac{1 - \theta_i(q_T)}{1 - \theta_i(q)} \geq \frac{1 - \theta_{i'}(q_T)}{1 - \theta_{i'}(q)}$ implies that $\frac{\theta_i(q_T) - \theta_i(q)}{1 - \theta_i(q)} \leq \frac{\theta_{i'}(q_T) - \theta_{i'}(q)}{1 - \theta_{i'}(q)}$, and so it follows from Lemma 1 that the above result is nonnegative for $i \leq i^*$ and nonpositive for $i > i^*$.

Repeating the above construction at all period T histories $h_T \in \mathcal{H}_T$ yields a test such that all deterministic tests with random verdicts in its support assign task q at period T and full effort is optimal for all types of the agent at all period T histories. Moreover, since this (weakly) raises the payoffs of good types and lowers those of bad types at all period T histories, it does not lower the principal's payoff.

We now proceed inductively backward from period $T - 1$. For a given period $1 \leq t \leq T - 1$, we assume as the induction hypothesis that it is optimal for all types of the agent to choose full effort at all histories $h_{t'} \in \mathcal{H}_{t'}$ for $t < t' \leq T$ in all deterministic tests with random verdicts $(\mathcal{T}, \mathcal{V})$ that are in the support of ρ . Additionally, we assume as part of the induction hypothesis that $\mathcal{T}(h_{t'}) = q$ at all $h_{t'} \in \mathcal{H}_{t'}$ for $t < t' \leq T$.

Now consider each period t history $h_t \in \mathcal{H}_t$ and assigned task q_t . A consequence of the induction hypothesis is that it is without loss to assume that each $(\mathcal{T}, \mathcal{V}) \in \text{supp}(h_t, q_t)$ (if nonempty) has the same verdict at each terminal history in $\Gamma(h_t)$. This follows because, as per the induction hypothesis, only task q is assigned in periods $t + 1$ onward in the subtree $\Lambda(h_t)$, and so the agent learns nothing further as the test progresses. In other words, it is equivalent to set the verdicts of each $(\mathcal{T}, \mathcal{V}) \in \text{supp}(h_t, q_t)$ to be $\mathcal{V}(h_{T+1}) = \mathbb{E}[\mathcal{V}'(h_{T+1}) | (\mathcal{T}', \mathcal{V}') \in \text{supp}(h_t, q_t)]$ for all $h_{T+1} \in \Gamma(h_t)$.

We now alter each $(\mathcal{T}, \mathcal{V}) \in \text{supp}(h_t, q_t)$ so that task q is assigned at h_t and we change the verdicts so that full effort is optimal for the agent at all histories in $\Lambda(h_t)$. First, observe that following the argument of Step 1 of Lemma 2, we can assume that the verdicts $\mathcal{V}$ at terminal histories $\Gamma(\{h_t, (q_t, s)\})$ and $\Gamma(\{h_t, (q_t, f)\})$ satisfy the cutoff property of Observation 2.

Recall that a consequence of the above argument (Step 2 of Lemma 2) is that all types have the same optimal action at h_t , since the same task q is assigned at all histories from $t + 1$ onward in the subtree $\Lambda(h_t)$ and the verdicts satisfy the cutoff property. If the agent finds it optimal to shirk at h_t , then we can construct $(\mathcal{T}', \mathcal{V}')$, which is identical to $(\mathcal{T}, \mathcal{V})$ except that the verdicts at terminal histories $\{h_t, (q_t, s), h'\} \in \Gamma(\{h_t, (q_t, s)\})$ are reassigned to those in $\Gamma(\{h_t, (q_t, f)\})$ by setting $\mathcal{V}'(\{h_t, (q_t, s), h'\}) = \mathcal{V}(\{h_t, (q_t, f), h'\})$. This would make all types indifferent among all actions and would not change their payoffs or the payoff of the principal. Moreover, this replacement of verdicts makes the task at h_t irrelevant, so that we can replace q_t with q at h_t (and reassign the verdicts accordingly).

Now consider the case in which action 1 is optimal for all types at h_t . We now replace each $(\mathcal{T}, \mathcal{V}) \in \text{supp}(h_t, q_t)$ by another test $(\mathcal{T}', \mathcal{V}')$. As in the argument for period T above, we consider two separate cases.

Let $\beta'_1 := \max_{i' \leq i^*} \frac{\theta_{i'}(q_t)}{\theta_{i'}(q)}$. First, suppose $\beta'_1 \leq 1$. Then we take the test $(\mathcal{T}', \mathcal{V}')$ to be identical to $(\mathcal{T}, \mathcal{V})$ except that $\mathcal{T}'(h_t) = q$ and the verdicts at the terminal histories $\{h_t, (q, s), h'\} \in \Gamma(\{h_t, (q, s)\})$ are $\mathcal{V}'(\{h_t, (q, s), h'\}) = \beta'_1 \mathcal{V}(\{h_t, (q_t, s), h'\}) + (1 - \beta'_1) \mathcal{V}(\{h_t, (q_t, f), h'\})$. In words, we are replacing the verdicts following a success at h_t with a weighted average of the verdicts following a success and failure before the change.

For brevity, we define

$$u_i^s := u_i(\{h_t, (q_t, s)\}; \mathcal{T}, \mathcal{V}, \sigma_i^*) \quad \text{and} \quad u_i^f := u_i(\{h_t, (q_t, f)\}; \mathcal{T}, \mathcal{V}, \sigma_i^*)$$

to be the expected payoffs following success and failure, respectively, at h_t in test $(\mathcal{T}, \mathcal{V})$.

We now show that this change (weakly) raises payoffs of good types and lowers those of bad types. Since full effort is optimal in the modified test, the payoff of type i at h_t from $(\mathcal{T}', \mathcal{V}')$ is

$$\theta_i(q)(\beta'_1 u_i^s + (1 - \beta'_1) u_i^f) + (1 - \theta_i(q)) u_i^f.$$

Following the same argument as for (8) (with u_i^s and u_i^f in place of v^s and v^f), the change in expected payoff at history h_t is given by

$$\theta_i(q_t)(u_i^s - u_i^f) \left(\frac{\theta_i(q)}{\theta_i(q_t)} \max_{i' \leq i^*} \left\{ \frac{\theta_{i'}(q_t)}{\theta_{i'}(q)} \right\} - 1 \right),$$

which is nonnegative for $i \leq i^*$ and nonpositive for $i > i^*$.

A similar construction can be used for the second case where $\beta'_1 > 1$. In this case, we take the test $(\mathcal{T}', \mathcal{V}')$ to be identical to $(\mathcal{T}, \mathcal{V})$ except that $\mathcal{T}'(h_t) = q$ and the verdicts at the terminal histories $\{h_t, (q, f), h'\} \in \Gamma(\{h_t, (q, f)\})$ are $\mathcal{V}'(\{h_t, (q, f), h'\}) = \beta'_2 \mathcal{V}(\{h_t, (q_t, f), h'\}) + (1 - \beta'_2) \mathcal{V}(\{h_t, (q_t, s), h'\})$, where $\beta'_2 := 1 - \max_{i' \leq i^*} \frac{\theta_{i'}(q_t) - \theta_{i'}(q)}{1 - \theta_{i'}(q)}$. In words, we are replacing the verdicts following a failure at h_t with a weighted average of the verdicts following a success and failure before the change.

As before, the difference in payoffs is

$$(\theta_i(q_t) - \theta_i(q))(u_i^s - u_i^f) \left(\frac{1 - \theta_i(q)}{\theta_i(q_t) - \theta_i(q)} \max_{i' \leq i^*} \frac{\theta_{i'}(q_t) - \theta_{i'}(q)}{1 - \theta_{i'}(q)} - 1 \right),$$

which is nonnegative for $i \leq i^*$ and nonpositive for $i > i^*$.

Repeating this construction at all period t histories completes the induction step and, therefore, also the proof. $\square$

Proof of Theorem 3

Suppose that $\pi_{\bar{i}} = \pi_{\bar{j}} = 0.5$. Let ρ be a test for which $\mathcal{T}(h) \equiv q$ for every $(\mathcal{T}, \mathcal{V})$ in the support of ρ . Since $\theta_{\bar{j}}(q) > \theta_{\bar{i}}(q)$ for any strategy of type $\bar{i}$, there exists a strategy of type $\bar{j}$ that generates the same distribution over terminal histories. In particular, it must be that $v_{\bar{j}}(\rho) \geq v_{\bar{i}}(\rho)$, which in turn implies that the principal's expected payoff is nonpositive.

Let q' be such that $\theta_i(q') = 1 - \theta_i(q)$ for every i . Notice that q is more Blackwell informative than q' since (3) is satisfied with $\alpha_s = 0$ and $\alpha_f = 1$.²³ Consider the test $(\mathcal{T}', \mathcal{V}')$ such that $\mathcal{T}'(h) \equiv q'$ and $\mathcal{V}'(h) = 1$ if and only if $h = ((q', s), \dots, (q', s))$; in words, the test always assigns q' and passes the agent if and only if he succeeds in every period. Given this test, the full-effort strategy is optimal for the agent, and $v_{\bar{i}}(\mathcal{T}', \mathcal{V}') > v_{\bar{j}}(\mathcal{T}', \mathcal{V}')$ since $\theta_{\bar{i}}(q') > \theta_{\bar{j}}(q')$. Therefore, the principal's expected payoff,

$$0.5v_{\bar{i}}(\mathcal{T}', \mathcal{V}') - 0.5v_{\bar{j}}(\mathcal{T}', \mathcal{V}'),$$

is positive, which in turn implies that this test is strictly better than any test that assigns q at every history. $\square$

²³The comparison between q and q' is weak in the sense that q' is also more Blackwell informative than q . An identical argument applies if instead q' solves (3) for some α_s and α_f satisfying $0 < \alpha_s < \alpha_f < 1$, in which case q is strictly more Blackwell informative than q' .

Proof of Theorem 4

We first show that the full-effort strategy σ^N is optimal for the agent in some optimal test. Then we show that σ^N is also optimal for the agent in the OFT.

We show the first part by contradiction. Suppose ρ is an optimal test where there is at least one history where full effort is not optimal for the agent. We proceed backward from period T , altering each deterministic test with random verdicts in the support of ρ in a way that both types find it optimal to choose full effort without reducing the payoff of the principal.

Consider first a period T history h_T together with an assigned task q_T . Let

$$v^\omega := \mathbb{E}[\mathcal{V}(\{h_T, (q_T, \omega)\}) | (\mathcal{T}, \mathcal{V}) \in \text{supp}(h_T, q_T)]$$

be the expected verdict following the outcome $\omega \in \{s, f\}$ taken with the respect to the set of possible deterministic tests with random verdicts that the agent could be facing.

Suppose that shirking is optimal for some type θ_i . Then it must be that $v^f \geq v^s$, which in turn implies that both types find shirking optimal and thereby get expected verdict v^f . Replacing each deterministic test with random verdicts $(\mathcal{T}, \mathcal{V}) \in \text{supp}(h_T, q_T)$ with another $(\mathcal{T}', \mathcal{V}')$ that is identical except that $\mathcal{V}'(\{h_T, (q_T, s)\}) = v^f$ does not alter the payoffs of the principal or the agent and makes full effort optimal at h_T .

We now proceed inductively backward from period $T - 1$. For a given period $1 \leq t \leq T - 1$, we assume as the induction hypothesis that it is optimal for all types of the agent to choose full effort at all histories $h_{t'} \in \mathcal{H}_{t'}$ for $t < t' \leq T$ in all deterministic tests with random verdicts $(\mathcal{T}, \mathcal{V})$ that are in the support of ρ .

Now consider each period t history $h_t \in \mathcal{H}_t$ and a task q_t such that full effort is not optimal for at least one type of the agent. If no such period t history exists, the induction step is complete. We now alter each $(\mathcal{T}, \mathcal{V}) \in \text{supp}(h_t, q_t)$ so that $a_t = 1$ is optimal for the agent at all histories in $\Lambda(h_t)$. We consider two separate cases:

- (i) Shirking is optimal for the good type, i.e., $\sigma_1^*(h_t) = 0$.
- (ii) Shirking is optimal for the bad type and full effort is optimal for the good type, i.e., $\sigma_2^*(h_t) = 0$ and $\sigma_1^*(h_t) = 1$.

In case (i), we replace each $(\mathcal{T}, \mathcal{V}) \in \text{supp}(h_t, q_t)$ by $(\mathcal{T}', \mathcal{V}')$ where the continuation test following the success is replaced by that following a failure. Formally, $(\mathcal{T}', \mathcal{V}')$ is identical to $(\mathcal{T}, \mathcal{V})$ except for the tasks and verdicts in the subtree $\Lambda(\{h_t, (q_t, s)\})$. For each history $\{h_t, (q_t, s), h'\} \in \Lambda(\{h_t, (q_t, s)\})$ in this subtree, the task assigned becomes $\mathcal{T}'(\{h_t, (q_t, s), h'\}) = \mathcal{T}(\{h_t, (q_t, f), h'\})$, and the verdict at each terminal history $\{h_t, (q_t, s), h'\} \in \Gamma(\{h_t, (q_t, s)\})$ becomes $\mathcal{V}'(\{h_t, (q_t, s), h'\}) = \mathcal{V}(\{h_t, (q_t, f), h'\})$. Note that if we alter each $(\mathcal{T}, \mathcal{V}) \in \text{supp}(h_t, q_t)$ in this way, the performance of the agent at h_t does not affect the expected verdict and so $a_t = 1$ is optimal for both types. By the induction hypothesis, action 1 remains optimal for both types at all histories in the subtree $\Lambda(h_t)$. Finally, such an alteration does not affect the payoff of the good type and weakly decreases the payoff of the bad type at h_t , and, therefore, weakly increases the principal's payoff.

In case (ii), we do the opposite and replace each $(\mathcal{T}, \mathcal{V}) \in \text{supp}(h_t, q_t)$ by $(\mathcal{T}', \mathcal{V}')$, where the continuation test following the failure is replaced by that following a success. Formally, $(\mathcal{T}', \mathcal{V}')$ is identical to $(\mathcal{T}, \mathcal{V})$ except for the tasks and verdicts in the subtree $\Lambda(\{h_t, (q_t, f)\})$. For each history $\{h_t, (q_t, f), h'\} \in \Lambda(\{h_t, (q_t, f)\})$ in this subtree, the task assigned becomes $\mathcal{T}'(\{h_t, (q_t, f), h'\}) = \mathcal{T}(\{h_t, (q_t, s), h'\})$, and the verdict at each terminal history $\{h_t, (q_t, f), h'\} \in \Gamma(\{h_t, (q_t, f)\})$ becomes $\mathcal{V}'(\{h_t, (q_t, f), h'\}) = \mathcal{V}(\{h_t, (q_t, s), h'\})$. Once again, the performance of the agent at h_t does not affect the expected verdict and so $a_t = 1$ is optimal for both types. By the induction hypothesis, action 1 remains optimal for both types at all histories in the subtree $\Lambda(h_t)$. Finally, such an alteration neither increases the payoff of the bad type nor decreases the payoff of the good type at h_t , and, therefore, weakly increases the principal's payoff. This completes the induction step.

Finally, we show that σ^N is optimal for the agent in the OFT $(\mathcal{T}^N, \mathcal{V}^N)$. We prove the result by induction on T . The base case is trivial since $\mathcal{V}^N(\{h_T, (\mathcal{T}^N(h_T), s)\}) \geq \mathcal{V}^N(\{h_T, (\mathcal{T}^N(h_T), f)\})$ for any history $h_T \in \mathcal{H}_T$, and so action 1 is optimal in the last period of the OFT (which is the only period when $T = 1$).

As the induction hypothesis, we assume that full effort is always optimal for the agent when faced with the OFT and when the length of the test is $T - 1$ or less. Thus, for the induction step, we need to argue that full effort is optimal for the agent in period 1 when the length of the test is T .

Accordingly, suppose the agent has a strict preference to shirk in period 1. We consider three separate cases:

- (i) The good type strictly prefers to shirk while full effort is optimal for the bad type; thus $\sigma_1^*(h_t) = 0$ and $\sigma_2^*(h_t) = 1$.
- (ii) The bad type strictly prefers to shirk while full effort is optimal for the good type; thus $\sigma_2^*(h_t) = 0$ and $\sigma_1^*(h_t) = 1$.
- (iii) Both types strictly prefer to shirk; thus $\sigma_1^*(h_t) = \sigma_2^*(h_t) = 0$.

Cases (i) and (ii) can be handled in the same way as cases (i) and (ii) from the first part of the proof. In case (i), the continuation test following a success is replaced by that following a failure. Given the strategy σ^N , this change strictly increases the payoff of the good type and weakly decreases the payoff of the bad type, contradicting the optimality of the OFT. For case (ii), the continuation test following the failure can be replaced by that following a success, providing the requisite contradiction.

Now consider case (iii). Let $h_2^s = \{(\mathcal{T}^N(h_1), s)\}$ and $h_2^f = \{(\mathcal{T}^N(h_1), f)\}$, and let $\pi_i^N(h)$ denote the belief the principal assigns to the agent's type being θ_i following history h under the assumption that the agent uses the full-effort strategy σ^N . Note that group monotonicity implies that $\pi_1(h_2^s) \geq \pi_1(h_2^f)$ (and, equivalently, $\pi_2(h_2^s) \leq \pi_2(h_2^f)$). If $\pi_1(h_2^s) = \pi_1(h_2^f)$, then it must be that there is no task q that satisfies $\theta_1(q) \neq \theta_2(q)$, for otherwise the OFT would assign such a task in the first period and $\pi_1(h_2^s)$ would differ from $\pi_1(h_2^f)$. In that case, the result holds trivially. Thus we may assume that $\pi_1(h_2^s) > \pi_1(h_2^f)$ and $\pi_2(h_2^s) < \pi_2(h_2^f)$.

By the optimality of the continuation test following a success, we have

$$\begin{aligned} & \pi_1^N(h_2^s)u_1(h_2^s; (\mathcal{T}^N, \mathcal{V}^N), \sigma_1^N) - \pi_2^N(h_2^s)u_2(h_2^s; (\mathcal{T}^N, \mathcal{V}^N), \sigma_2^N) \\ & \geq \pi_1^N(h_2^s)u_1(h_2^f; (\mathcal{T}^N, \mathcal{V}^N), \sigma_1^N) - \pi_2^N(h_2^s)u_2(h_2^f; (\mathcal{T}^N, \mathcal{V}^N), \sigma_2^N), \end{aligned}$$

since otherwise the principal would be better off replacing the continuation test after a success with that after a failure. Rearranging gives

$$\begin{aligned} & \pi_2^N(h_2^s)[u_2(h_2^f; (\mathcal{T}^N, \mathcal{V}^N), \sigma_2^N) - u_2(h_2^s; (\mathcal{T}^N, \mathcal{V}^N), \sigma_2^N)] \\ & \geq \pi_1^N(h_2^s)[u_1(h_2^f; (\mathcal{T}^N, \mathcal{V}^N), \sigma_1^N) - u_1(h_2^s; (\mathcal{T}^N, \mathcal{V}^N), \sigma_1^N)]. \end{aligned}$$

Similarly, by the optimality of the continuation test following a failure, we have

$$\begin{aligned} & \pi_1^N(h_2^f)[u_1(h_2^f; (\mathcal{T}^N, \mathcal{V}^N), \sigma_1^N) - u_1(h_2^s; (\mathcal{T}^N, \mathcal{V}^N), \sigma_1^N)] \\ & \geq \pi_2^N(h_2^f)[u_2(h_2^f; (\mathcal{T}^N, \mathcal{V}^N), \sigma_2^N) - u_2(h_2^s; (\mathcal{T}^N, \mathcal{V}^N), \sigma_2^N)]. \end{aligned}$$

Since $\pi_1^N(h_2^s) > \pi_1^N(h_2^f)$ and $u_1(h_2^f; (\mathcal{T}^N, \mathcal{V}^N), \sigma_1^N) > u_1(h_2^s; (\mathcal{T}^N, \mathcal{V}^N), \sigma_1^N)$ (since type θ_1 strictly prefers to shirk), the above two inequalities imply that

$$\begin{aligned} & \pi_2^N(h_2^s)[u_2(h_2^f; (\mathcal{T}^N, \mathcal{V}^N), \sigma_2^N) - u_2(h_2^s; (\mathcal{T}^N, \mathcal{V}^N), \sigma_2^N)] \\ & \geq \pi_2^N(h_2^f)[u_2(h_2^f; (\mathcal{T}^N, \mathcal{V}^N), \sigma_2^N) - u_2(h_2^s; (\mathcal{T}^N, \mathcal{V}^N), \sigma_2^N)]. \end{aligned}$$

Since $u_2(h_2^f; (\mathcal{T}^N, \mathcal{V}^N), \sigma_2^N) > u_2(h_2^s; (\mathcal{T}^N, \mathcal{V}^N), \sigma_2^N)$ (since type θ_2 also strictly prefers to shirk), this inequality implies that $\pi_2^N(h_2^s) \geq \pi_2^N(h_2^f)$, a contradiction. $\square$

APPENDIX B: ADDITIONAL EXAMPLES

EXAMPLE 7. This example demonstrates that (i) strategic behavior by the agent can be harmful to the principal and yield her a lower payoff than when the agent chooses σ^N in the OFT, and (ii) the optimal deterministic test may differ from the OFT even if σ^N is optimal for the agent in the former (but not in the latter).

Suppose there are three types ($I = 3$) and three periods ($T = 3$), with $i^* = 1$ (so that type θ_1 is the only good type). The principal has two different tasks, $Q = \{q, q'\}$, and the success probabilities are

	q	q'
θ_1	1	0.9
θ_2	0.85	0.8
θ_3	0.8	0.

The principal's prior belief is

$$(\pi_1, \pi_2, \pi_3) = (0.4, 0.1, 0.5).$$

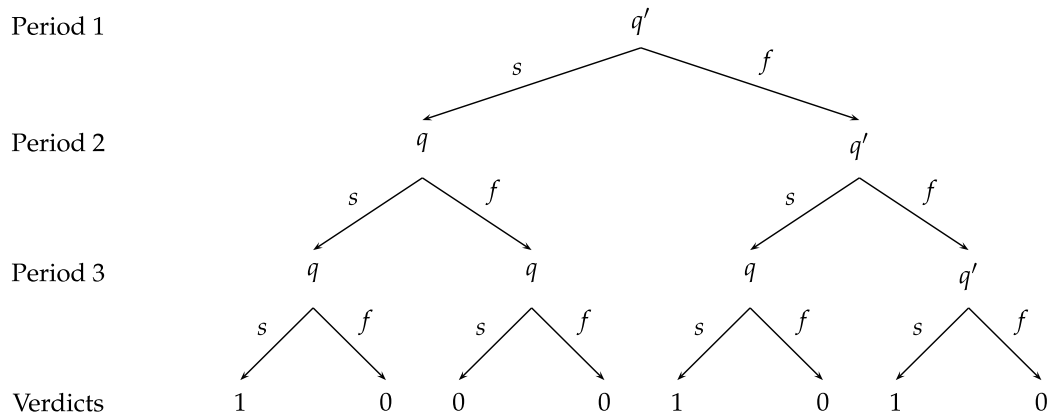


FIGURE 5. An OFT for Example 7.

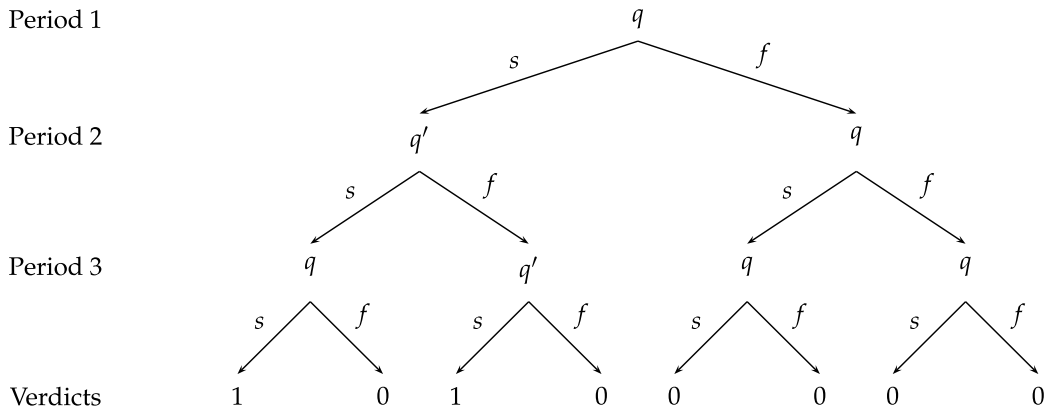


FIGURE 6. An optimal deterministic test for Example 7.

Figure 5 depicts an OFT $(\mathcal{T}^N, \mathcal{V}^N)$. The intuition for this OFT is as follows. The prior probability is such that type θ_2 is unlikely, and task q' is more effective at differentiating between types θ_1 and θ_3 . However, type θ_3 never succeeds at task q' , so as soon as a success is observed, the principal concludes that the agent's type must be either θ_1 or θ_2 and switches to task q (which is better at differentiating between these types).

Note that full effort is not optimal for the agent in this test: type θ_2 prefers to choose effort 0 in period 1 because his expected payoff $u_2(h_2; \mathcal{T}^N, \mathcal{V}^N, \sigma_2^N) = 0.85 * 0.85 = 0.7225$ at history $h_2 = \{(q', s)\}$ is lower than $u_2(h'_2; \mathcal{T}^N, \mathcal{V}^N, \sigma_2^N) = 0.8 * 0.85 + 0.2 * 0.8 = 0.84$ at the history $h'_2 = \{(q', f)\}$. This shirking lowers the principal's payoff since it increases the payoff of a bad type.

An optimal deterministic test $(\mathcal{T}', \mathcal{V}')$ is depicted in Figure 6. Observe that in this test, σ^N is an optimal strategy for the agent. Here, the principal screens in period 1 by assigning task q instead of q' , following which she assigns the two-period OFT (for the corresponding posterior beliefs). Note that the posterior belief following a failure on the period 1 task q assigns zero probability to the agent being type θ_1 , which implies

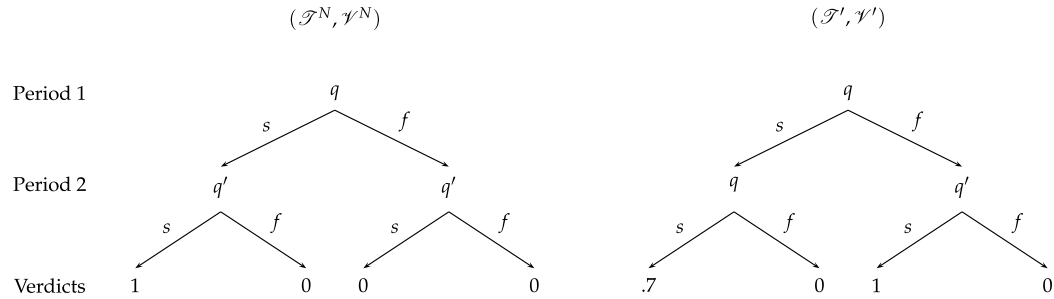


FIGURE 7. Tests for Example 8. The test on the left is an OFT and that on the right is an optimal test.

the agent will never pass the test. Following a success, the same intuition from above applies: type θ_2 is less likely than types θ_1 and θ_3 , and so the principal assigns task q' at history $h_2 = \{(q, s)\}$ and switches to q only if the agent succeeds at this task (revealing that he is not θ_3). By definition, since the agent chooses the full-effort strategy, this test must yield a lower payoff to the principal than she would obtain if the agent chose σ^N in the OFT. $\diamond$

EXAMPLE 8. The main purpose of this example is to demonstrate that the optimal test may employ a less informative task even if group monotonicity holds. In other words, Theorem 1 cannot be strengthened to state that less informative tasks are not used in the optimal test when there does not exist a single most informative task. This example also shows that the principal can sometimes benefit from randomization: the optimal deterministic test in this case gives the principal a lower payoff than does the optimal test. This benefit arises from randomizing verdicts, but in a variant of this example, the principal can do strictly better by randomizing tasks.

This example features three types ($I = 3$) and two periods ($T = 2$), with $i^* = 2$ (so that type θ_3 is the only bad type). Suppose first that the principal has two different tasks, $Q = \{q, q'\}$, with the success probabilities

	q	q'
θ_1	0.9	0.5
θ_2	0.4	0.35
θ_3	0.3	0.21.

The principal's prior belief is

$$(\pi_1, \pi_2, \pi_3) = (0.02, 0.4, 0.58).$$

Figure 7 depicts, on the left, an OFT $(\mathcal{T}^N, \mathcal{V}^N)$ (which is also an optimal deterministic test), and, on the right, an optimal test $(\mathcal{T}', \mathcal{V}')$. The test $(\mathcal{T}', \mathcal{V}')$ differs from $(\mathcal{T}^N, \mathcal{V}^N)$ in two ways: task q' at history $\{(q, s)\}$ is replaced by task q , and the verdicts at both terminal histories involving a success in period 2 are changed. Note that, in period 1, types θ_1 and θ_2 strictly prefer actions $a_1 = 1$ and $a_1 = 0$, respectively, whereas type θ_3 is indifferent.

The following simple calculations demonstrate why $(\mathcal{T}', \mathcal{V}')$ yields the principal a higher payoff than does $(\mathcal{T}^N, \mathcal{V}^N)$. In $(\mathcal{T}', \mathcal{V}')$, the payoff of all three types is higher than in $(\mathcal{T}^N, \mathcal{V}^N)$. The differences in payoffs are

$$\Delta v_1 = v_1(\mathcal{T}', \mathcal{V}') - v_1(\mathcal{T}^N, \mathcal{V}^N) = 0.9 * 0.9 * 0.7 + 0.1 * 0.5 - 0.9 * 0.5 = 0.167$$

$$\Delta v_2 = v_2(\mathcal{T}', \mathcal{V}') - v_2(\mathcal{T}^N, \mathcal{V}^N) = 0.35 - 0.4 * 0.35 = 0.21$$

$$\Delta v_3 = v_3(\mathcal{T}', \mathcal{V}') - v_3(\mathcal{T}^N, \mathcal{V}^N) = 0.21 - 0.3 * 0.21 = 0.147.$$

The change in the principal's payoff is

$$\sum_{i=1}^2 \pi_i \Delta v_i - \pi_3 \Delta u_3 = 0.02 * 0.167 + 0.4 * 0.21 - 0.58 * 0.147 > 0,$$

which implies that $(\mathcal{T}', \mathcal{V}')$ is better than $(\mathcal{T}^N, \mathcal{V}^N)$ for the principal.

Proving that $(\mathcal{T}', \mathcal{V}')$ is optimal is more challenging; we provide a sketch of the argument here. Whenever there is a single bad type, there is an optimal test that satisfies at least one of the following two properties: (i) there is no randomization of tasks in period two or (ii) the bad type is indifferent among all actions in period 1. To see this, suppose, to the contrary, that the bad type has a strictly optimal action in period 1, and that the principal assigns probability $\beta \in (0, 1)$ to q and $1 - \beta$ to q' at one of the histories in period 2. Observe that for a fixed strategy of the agent, the principal's payoff is linear in this probability β . Hence the principal can adjust β without lowering her payoff until either θ_3 becomes indifferent in period 1 or β becomes 0 or 1; any change in the strategies of types θ_1 and θ_2 resulting from this adjustment only benefits the principal more. Establishing that the optimal test must satisfy (i) or (ii) makes it possible to show that $(\mathcal{T}', \mathcal{V}')$ is optimal by comparing the principal's payoffs from tests having one of these properties.

Now suppose the principal has at her disposal another task q'' that satisfies

$$\theta_i(q'') = \theta_i(q) + \alpha(1 - \theta_i(q))$$

for all $i \in \{1, 2, 3\}$ and some $\alpha \in (0, 1]$. Task q is more Blackwell informative than q'' (one can take $\alpha_s = 1$ and $\alpha_f = \alpha$ in (3)).

The principal can now increase her payoff relative to $(\mathcal{T}', \mathcal{V}')$ by using the less informative task q'' . To see this, suppose the principal assigns q'' instead of q in the first period, without changing tasks and verdicts in period 2. This change does not affect the payoffs or optimal strategies of types θ_2 and θ_3 ; the former still chooses $a_1 = 0$ and the latter remains indifferent among all actions. However, this change does increase the payoff of type θ_1 , since this type strictly prefers the subtree after a success in period 1 to that after a failure, and task q'' gives a higher probability of reaching this subtree than does q . Therefore, this change increases the principal's payoff and demonstrates that any optimal test with the set of tasks $\{q, q', q''\}$ must employ q'' .

Finally, to show that the principal can sometimes benefit from randomizing tasks, suppose that the set of tasks is given by $\{q, q', q'''\}$, where $\theta_1(q''') = 0.5$, $\theta_2(q''') = 0.16$,

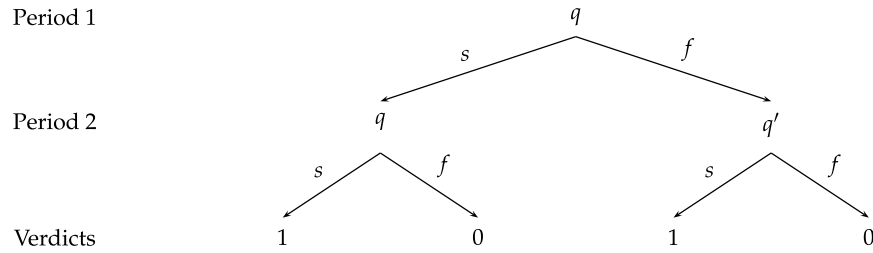


FIGURE 8. The optimal test for Example 9.

and $\theta_3 = 0.12$. Consider the test that assigns task q in the first period, and in the second period assigns q' if the agent failed on the first task while randomizing equally between q and q'' if the agent succeeded in the first period. The verdict passes the agent if and only if he succeeds on the task in period 2. For this test, the probabilities of passing for types θ_2 and θ_3 are identical to those in the test on the right-hand side of Figure 7; the only difference is that type θ_1 is more likely to pass the test. By checking various cases, one can show that the optimal test that does not randomize tasks never assigns q'' . Therefore, the principal strictly benefits from randomizing tasks. $\diamond$

EXAMPLE 9. This example extends Example 5 to show that the principal can benefit from offering a menu of tests. Recall that the success probabilities are

	q	q'
θ_1	1	0.2
θ_2	0.2	0.15
θ_3	0.1	0.01

and the prior is

$$(\pi_1, \pi_2, \pi_3) = (0.5, 0.1, 0.4).$$

Suppose that there are only two periods ($T = 2$).

The test depicted in Figure 8 is the optimal deterministic test (and also the OFT). Observe that in this test, a failure in period 1 results in a harder task and that a success in period 2 is required to pass. Types θ_1 , θ_2 , and θ_3 pass with probabilities 1, $0.2 * 0.2 + 0.8 * 0.15 = 0.16$, and $0.1 * 0.1 + 0.9 * 0.01 = 0.019$, respectively.

Now suppose the principal instead offers the two-test menu $\{(\mathcal{T}_1, \mathcal{V}_1), (\mathcal{T}_2, \mathcal{V}_2)\}$ depicted in Figure 9. Note that the test $(\mathcal{T}_1, \mathcal{V}_1)$ only assigns the easier task, q , and two successes are required to pass. In contrast, test $(\mathcal{T}_2, \mathcal{V}_2)$ assigns only the harder task, q' , but a single success in either period is sufficient to pass. It is optimal for type θ_1 to choose $(\mathcal{T}_1, \mathcal{V}_1)$ and then use the full-effort strategy, as doing so enables him to pass with probability 1. Types θ_2 and θ_3 prefer to choose $(\mathcal{T}_2, \mathcal{V}_2)$ and then use the full-effort strategy. For types θ_2 and θ_3 , the passing probabilities are $0.2 * 0.2 = 0.04$ and $0.1 * 0.1 = 0.01$, respectively, in test $(\mathcal{T}_1, \mathcal{V}_1)$, which are lower than the corresponding passing probabilities $0.15 + 0.85 * 0.15 = 0.2775$ and $0.01 + 0.99 * 0.01 = 0.0199$ in test $(\mathcal{T}_2, \mathcal{V}_2)$.

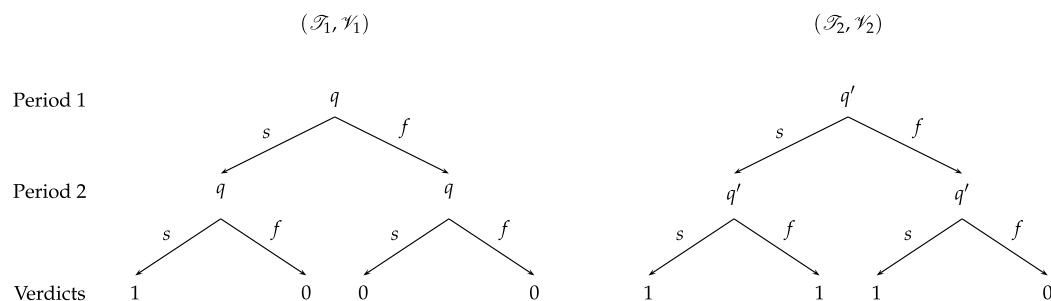


FIGURE 9. Menu of tests for Example 9.

Note that in this menu, the payoffs of types θ_2 and θ_3 go up relative to what they obtain in the optimal test. However, the gain for type θ_2 is much larger than for θ_3 , making the principal better off overall. In other words, the menu strictly increases the principal's payoff above that from the optimal test. $\diamond$

EXAMPLE 10. This example shows that if the principal cannot commit, she may not be able to implement the optimal test. Consider the following minor modification of the success probabilities from Example 4:

	q	q'
θ_1	0.999	0.5
θ_2	0.5	0.5
θ_3	0.5	0.4.

Note that the only change is that we have replaced $\theta_1(q) = 1$ by $\theta_1(q) = 0.999$. The prior remains unchanged. Since the payoffs are continuous in these probabilities, this minor modification affects neither the OFT nor the optimal test.

Suppose the optimal test could be implemented without commitment. Recall that type θ_1 chooses the full-effort strategy, whereas types θ_2 and θ_3 choose $a_t = 0$ in periods 1 and 2. This implies that the terminal histories $\{(q, s), (q, f), (q', s)\}$ and $\{(q, s), (q, f), (q', f)\}$ are never reached by θ_2 and θ_3 in equilibrium. However, there is a positive (albeit small) probability that these terminal histories are reached by type θ_1 . Therefore, a sequentially rational principal would assign verdicts 1 (instead of 0) at both of these terminal histories, which would in turn make full effort optimal for types θ_2 and θ_3 in the first period. $\diamond$

REFERENCES

- Ashworth, Scott, E. de Mesquita, and Amanda Friedenberg (2017), "Accountability and information in elections." *American Economic Journal: Microeconomics*, 9, 95–138. [1237]
- Baker, George P., Michael C. Jensen, and Kevin J. Murphy (1988), "Compensation and incentives: Practice vs. theory." *Journal of Finance*, 43, 593–616. [1235]

- Bergemann, Dirk. and J. Välimäki (2006), “Bandit problems.” Unpublished paper, Cowles Foundation Discussion Paper 1551. [1237]
- Blackwell, David (1953), “Equivalent comparisons of experiments.” *Annals of Mathematical Statistics*, 24, 265–272. [1234, 1242]
- Bradt, Russell N. and Samuel Karlin (1956), “On the design and comparison of certain dichotomous experiments.” *Annals of Mathematical Statistics*, 27, 390–409. [1241]
- DeGroot, Morris H. (1962), “Uncertainty, information, and sequential experiments.” *Annals of Mathematical Statistics*, 33, 404–419. [1234]
- DeGroot, Morris H. (2005), *Optimal Statistical Decisions*, volume 82 of Wiley Classics Library. John Wiley & Sons, Hoboken, New Jersey. [1234, 1242]
- Dewatripont, Mathias, Ian Jewitt, and Jean Tirole (1999), “The economics of career concerns: Part I.” *Review of Economic Studies*, 66, 183–198. [1236]
- Ferreira, Daniel and Raaj K. Sah (2012), “Who gets to the top? Generalists versus specialists in managerial organizations.” *The RAND Journal of Economics*, 43, 577–601. [1235]
- Foster, Dean and Rakesh Vohra (2011), “Calibration: Respice, adspice, prospice.” In *Advances in Economics and Econometrics: Tenth World Congress*, volume 1 (Daron Acemoglu, Manuel Arellano, and Eddie Dekel, eds.), 423–442, Cambridge University Press. [1237]
- Gerardi, Dino and Lucas Maestri (2012), “A principal–agent model of sequential testing.” *Theoretical Economics*, 7, 425–463. [1237]
- Gershkov, Alex and Motty Perry (2012), “Dynamic contracts with moral hazard and adverse selection.” *Review of Economic Studies*, 79, 268–306. [1237]
- Guasch, J. Luis and Andrew Weiss (1980), “Wages as sorting mechanisms in competitive markets with asymmetric information: A theory of testing.” *Review of Economic Studies*, 47, 653–664. [1236]
- Halac, Marina, Navin Kartik, and Qingmin Liu (2016), “Optimal contracts for experimentation.” *Review of Economic Studies*, 83, 1040–1091. [1237]
- Holmström, Bengt (1999), “Managerial incentive problems: A dynamic perspective.” *The Review of Economic Studies*, 66, 169–182. [1236]
- Hölmstrom, Bengt and Paul Milgrom (1987), “Aggregation and linearity in the provision of intertemporal incentives.” *Econometrica*, 55, 303–328. [1235, 1237, 1245]
- Laffont, Jean-Jacques and Jean Tirole (1988), “The dynamics of incentive contracts.” *Econometrica*, 56, 1153–1175. [1250]
- Meyer, Margaret A. (1994), “The dynamics of learning with team production: Implications for task assignment.” *Quarterly Journal of Economics*, 109, 1157–1184. [1234, 1240]

Meyer, Margaret A. and John Vickers (1997), “Performance comparisons and dynamic incentives.” *Journal of Political Economy*, 105, 547–581. [1250]

Milgrom, Paul R. and John Donald Roberts (1992), *Economics, Organization and Management*. Prentice-Hall. [1236, 1250]

Mirman, Leonard J., Larry Samuelson, and Edward E. Schlee (1994), “Strategic information manipulation in duopolies.” *Journal of Economic Theory*, 62, 363–384. [1237]

Nalebuff, Barry and David Scharfstein (1987), “Testing in models of asymmetric information.” *Review of Economic Studies*, 54, 265–277. [1236]

Olszewski, Wojciech (2015), “Calibration and expert testing.” In *Handbook of Game Theory*, volume 4 (H. Peyton Young and Shmuel Zamir, eds.), 949–984, North-Holland. [1237]

Prasad, Suraj (2009), “Task assignments and incentives: Generalists versus specialists.” *RAND Journal of Economics*, 40, 380–403. [1235]

Prescott, Edward C. and Michael Visscher (1980), “Organization capital.” *Journal of Political Economy*, 88, 446–461. [1233]

Rogerson, William P. (1985), “Repeated moral hazard.” *Econometrica*, 53, 69–76. [1237]

Co-editor George J. Mailath handled this manuscript.

Manuscript received 9 June, 2017; final version accepted 20 December, 2017; available online 27 December, 2017.