

Dynamic contracting: An irrelevance theorem

PÉTER ESŐ

Department of Economics, Oxford University

BALÁZS SZENTES

Department of Economics, London School of Economics

This paper generalizes a conceptual insight in dynamic contracting with quasilinear payoffs: the principal does not need to pay any information rents for extracting the agent's "new" private information obtained after signing the contract. This is shown in a general model in which the agent's type stochastically evolves over time, and her payoff (which is linear in transfers) depends on the entire history of private and any contractible information, contractible decisions, and her hidden actions. The contract is offered by the principal in the presence of initial informational asymmetry. The model can be transformed into an equivalent one where the agent's subsequent information is independent in each period (type orthogonalization). We show that for any fixed decision–action rule implemented by a mechanism, the agent's rents (as well as the principal's maximal revenue) are the same *as if* the principal could observe and contract on the agent's orthogonalized types after the initial period. We also show that any monotonic decision–action rule can be implemented in a Markovian environment satisfying certain regularity conditions, and we provide a simple "recipe" for solving such dynamic contracting problems.

KEYWORDS. Asymmetric information, dynamic contracting, mechanism design.

JEL CLASSIFICATION. D82, D83, D86.

1. INTRODUCTION

Two of the fundamental questions in contract theory and mechanism design are (i) what determines an agent's *rents* (or the uninformed principal's *agency costs*) in a given, implementable allocation rule or in the optimal policy, and, a related matter, (ii) what are the *optimal contracts* given the principal-designer's objective. In this paper we provide answers to both types of questions in a general *dynamic contracting environment* with quasilinear payoffs. Our multiperiod principal–agent model involves a stochastically evolving hidden type as well as hidden actions on the agent's part. The contractible per-period decisions and monetary transfers are governed by a contract signed after the

Péter Eső: peter.eso@economics.ox.ac.uk

Balázs Szentes: b.szentes@lse.ac.uk

We thank a co-editor and two anonymous referees for valuable comments. The paper also benefited from feedback of seminar participants at LSE, Oxford, and Yale.

Copyright © 2017 The Authors. Theoretical Economics. The Econometric Society. Licensed under the Creative Commons Attribution-NonCommercial License 3.0. Available at <http://econtheory.org>. DOI: 10.3982/TE2127

agent has learned her initial type. The agent's final payoff can depend, quite generally, on the histories of her private and (any contractible) public information, her hidden actions, and the contractible decisions, and it is linear in transfers between the parties. A wide range of applications of such models is discussed below.

Our main result is to show that the agent's rents (as well as the principal's maximal payoff, if it is also quasilinear) in a given, implementable allocation rule are the same *as if* the principal could observe and contract on any "new" (orthogonal) information observed by the agent after the contract is signed. In the hypothetical benchmark case where the agent's future, orthogonalized types are observable and contractible, the principal does not need the agent to report any new information beyond the initial period. Therefore, as far as the expected transfers implementing the allocation rule are concerned, it is inconsequential that in the original problem there is dynamic interaction between the parties. We term this result a *dynamic irrelevance theorem*. It holds in a rich environment, with very little assumed about the agent's utility function (no single-crossing or monotonicity assumptions are made), the information structure, and so on.

The dynamic irrelevance (or payment-characterization) result suggests a simple "recipe" for solving dynamic contracting problems in which the principal's payoff is quasilinear in payments. First, solve the *benchmark case* (where the agent's only private information is her initial type and all her orthogonalized future types are observable by the principal), which is essentially a static problem, hence standard solution methods are applicable. Second, check whether the solution to the benchmark (a decision–action rule) is *implementable* in the original model where the agent's type history is privately known. If it is, then the solution has been found. Note that the irrelevance result, which applies to a given, implementable decision–action rule, has no bearing on *whether* the benchmark solution is implementable in the original problem. Indeed, the two problems are *not equivalent*: the set of implementable decision–action rules is generally larger in the benchmark. To address the issue, our final set of results provides *sufficient conditions* for a decision–action rule to be implementable in the original problem.

The implementation results are derived in an environment where the agent's type follows a Markov process: her payoff is time-separable and it satisfies additional regularity (e.g., single-crossing) conditions. If in each period the principal observes a contractible signal that is informative (however imperfectly) about a summary statistic of the agent's type and action (e.g., the principal's profit that depends on the agent's quality and effort), and its distribution is generic, then *any monotonic decision rule* coupled with *any monotonic action rule* is approximately implementable.¹ If there are no such signals, then any *monotonic decision rule* coupled with *agent-optimal actions* is implementable. Therefore, after having solved the benchmark problem in a Markovian, regular environment, we only need to check if the solution satisfies the appropriate monotonicity condition; if it does, then the original problem is solved. In [Section 5](#) we present applications in which this approach can be used successfully.²

¹The genericity condition and the notion of implementability will be defined precisely in [Section 4](#).

²Little is known about conditions on the model's primitives under which implementability is the same in the original and benchmark problems. [Battaglini and Lamba \(2014\)](#) point out that the conditions for monotonicity of the "pointwise-optimal" decision rule can be quite strong.

One of the specific applications that we discuss is a *dynamic model of investment advice*. The principal is an investor (e.g., wealthy institution) and the agent is an advisor (e.g., private banker). The contractible variable is the amount of money that the principal invests with the agent in each period. The return on the investment is determined by the agent's type (evolving according to a first-order autoregressive process) and her costly, hidden effort in each period.³ The investor can observe and contract on a noisy public signal of the per-period return (e.g., a perceived return that may be affected by transitory, random events). The contract is offered by the investor after the advisor's initial type is realized. In this model we show that the advisor's rent (the investor's cost of agency) in a given, implementable decision–action rule only depends on her initial type realization (*dynamic irrelevance*). Any monotonic decision–action rule is approximately implementable. We derive the optimal contract, which is indeed monotonic and such that distortions dissipate over the long run. Besides this novel application we solve another two that are more familiar. The first one is a canonical dynamic monopoly problem in which the buyer's valuation for the good (her type) stochastically evolves over time. The second application differs from this as it allows the buyer to invest in her valuation by taking a private, costly action. We derive the optimal dynamic screening contract in the absence of any signal about the buyer's type and action. All distortions are due to the buyer's initial private information, again illustrating our dynamic irrelevance result.

Models in the class of dynamic contracting problems that we analyze have already been applied to a wide range of economic problems.⁴ The roots of this literature reach back to [Baron and Besanko \(1984\)](#) who used a multiperiod screening model to address the issue of regulating a monopoly over time. [Courty and Li \(2000\)](#) studied optimal advance ticket sales, [Eső and Szentes \(2007a\)](#) studied the optimal disclosure of private information in auctions, and [Eső and Szentes \(2007b\)](#) studied the sale of advice as an experience good. [Farhi and Werning \(2013\)](#), [Golosov et al. \(2011\)](#), and [Kapička \(2013\)](#) apply a similar approach to optimal taxation and fiscal policy design, respectively. [Pavan et al. \(2014\)](#) apply their (to date, most general) results on the multiperiod pure adverse selection problem to the auction of experience goods (bandit auctions). [Garrett and Pavan \(2012\)](#) use a dynamic contracting model with both adverse selection and moral hazard to study optimal CEO (chief operating officer) compensation. Such mixed, hidden-action–hidden-information models could also be applied in insurance problems.

Compared to the received literature (e.g., see the recent work of [Pavan et al. 2014](#)), we not only generalize the model to accommodate hidden actions and contractible signals, but we also state the main result as one of *irrelevance*: In these dynamic problems it is inconsequential whether the agent has access to dynamic deviation strategies. This result is obtained using an *orthogonalized representation* of the agent's private information: the agent's type in each period is constructed to be independent conditional on the history of types, actions, and decisions (a transformation proposed in [Eső and](#)

³The “return” may be more broadly interpreted as a composite score of the investment's monetary and nonmonetary (e.g., ethical) gains, differentially affected by the agent's type and effort.

⁴Our review of applications is deliberately incomplete; for a more in-depth survey of this literature, see [Krähmer and Strausz \(2015a\)](#) or [Pavan et al. \(2014\)](#).

Szentes (2007a) for an independent private value (IPV) auction environment). We show that in the original problem, where the agent's orthogonalized future types and actions are *not* observable, in any incentive compatible mechanism, the agent's expected payoff conditional on her initial type is fully determined by her on-path (in the future, truthful) behavior. The validity of this envelope theorem-type argument rests on the conditional independence of future information. Therefore, the agent's expected payoff (and payments) coincide with those in the benchmark case, where the orthogonalized future types are publicly observable.

The results on the implementability of monotonic decision–action rules in regular, Markovian environments are first established *without hidden actions*. Formally, this special case (unlike the general result) follows from results in Pavan et al. (2014); the proof relies on showing that if the agent is untruthful in a given period in an incentive compatible mechanism, she immediately *undoes her lie* in the following period to make the principal's inferences in all future periods correct, and this pins down all continuation payoffs and allows induction on the number of periods. The new results for models with *both hidden information and hidden action* are obtained by appropriately reducing the general model to ones with only adverse selection.⁵ When there is no contractible signal about the agent's type and action, we show that the agent's per-period utility given her type and an agent-optimal action satisfies the conditions that apply in the model with pure adverse selection; hence any monotonic decision rule coupled with an agent-optimal action is implementable. When there is an imperfect contractible signal about a summary statistic of the agent's type and action, we consider a felicity function defined by the agent's per-period utility as if the summary statistic were contractible and the agent could be compelled to generate the contracted summary statistic consistent with her type report. We show that this felicity function satisfies the sufficient conditions applicable under pure adverse selection; hence any decision rule coupled with an action rule such that both are monotonic in the type (which is implied by monotonic decision–action rules) is implementable. Finally, we show that with only an imperfect signal about the summary statistic, any monotonic decision–action rule may be implemented *in approximation*.

The technical contributions notwithstanding, we believe the most important message of the paper is the dynamic irrelevance result. The insight that the principal need not pay his agent rents for post-contractual hidden information in dynamic adverse selection has been expressed in previous work going back to Baron and Besanko (1984). Our paper highlights both the depth and the limitations of this insight: Indeed the principal who contracts the agent prior to her discovery of new information can limit the agent's rents as if he could observe the agent's orthogonalized future types; however, the two problems are not equivalent.

The paper is organized as follows. In Section 2 we introduce the model and describe the orthogonal transformation of the agent's information. In Section 3 we derive necessary conditions of the implementability of a decision rule and our main, dynamic irrelevance result. Section 4 presents sufficient conditions for implementation in Markov

⁵This technique is familiar from Laffont and Tirole (1986) and has also been used in a dynamic setting by Garrett and Pavan (2012).

environments. [Section 5](#) presents the applications; [Section 6](#) concludes. Omitted proofs are given in Appendices A and B. (Appendix B is available in a supplementary file on the journal website, <http://econtheory.org/supp/2127/supplement.pdf>.)

2. MODEL

Environment

There is a single principal (he) and a single agent (she). Time is discrete, indexed by $t = 0, 1, \dots, T < \infty$. The agent's private information in period t is $\theta_t \in \Theta_t$, where $\Theta_t = [\underline{\theta}_t, \bar{\theta}_t] \subset \mathbb{R}$. In period t the agent takes action $a_t \in A_t$, which is not observed by the principal. The set A_t is an open interval of $\mathbb{R}$.⁶ Then a contractible signal is drawn, $s_t \in S_t \subset \mathbb{R}$, that is observed by the principal. After s_t is realized in period t , a contractible decision is made, denoted by $x_t \in X_t \subset \mathbb{R}^n$. Since x_t is contractible, it does not matter whether it is taken by the agent or by the principal. The contract between the principal and the agent is signed at $t = 0$, right after the agent has learned her initial type, θ_0 .

We denote the history of a variable through period t by superscript t ; for example, $x^t = (x_0, \dots, x_t)$ and $x^{-1} = \{\emptyset\}$. The random variable θ_t is distributed according to a cumulative distribution function (c.d.f.) $G_t(\cdot | \theta^{t-1}, a^{t-1}, x^{t-1})$ supported on Θ_t . The function G_t is continuously differentiable in all of its argument, and the density is denoted by $g_t(\cdot | \theta^{t-1}, a^{t-1}, x^{t-1})$.

Signal s_t is distributed according to a continuous c.d.f. $\mathcal{H}_t(\cdot | \theta_t, a_t)$. We assume that when generating this signal, the agent is able to *compensate* her type by her action locally. Formally, there is a $\delta > 0$ such that for all $\hat{\theta}_t$, θ_t , and a_t , if $|\hat{\theta}_t - \theta_t| < \delta$, then there is an $\hat{a}_t$ such that $\mathcal{H}_t(\cdot | \theta_t, a_t) = \mathcal{H}_t(\cdot | \hat{\theta}_t, \hat{a}_t)$. This assumption ensures that the principal cannot resolve the adverse selection problem by requiring the agent to take a certain action and using signal s_t to detect the agent's type. A consequence of this assumption is that there exists a function $f_t : \Theta_t \times A_t \rightarrow \mathbb{R}$ such that the distribution of s_t depends only on $f_t(\theta_t, a_t)$, that is, $\mathcal{H}_t(\cdot | \theta_t, a_t) = H_t(\cdot | f_t(\theta_t, a_t))$ for some conditional c.d.f. H_t .⁷ We assume that f_t is continuously differentiable. We may interpret s_t as an imperfect public summary signal about the agent's current type and action; for example, in [Application 3](#) in [Section 5](#) it will be $s_t = \theta_t + a_t + \xi_t$, where ξ_t is noise with a known distribution.

The agent's total payoff is quasilinear in money and is defined by

$$\tilde{u}(\theta^T, a^T, s^T, x^T) - p,$$

where $p \in \mathbb{R}$ denotes the agent's payment to the principal, and $\tilde{u} : \Theta^T \times A^T \times S^T \times X^T \rightarrow \mathbb{R}$ is continuously differentiable in θ_t and a_t for all $t = 0, \dots, T$. We do not specify the principal's payoff. In some applications (e.g., where the principal is a monopoly and the agent is its customer) it could be the payment itself; in others (e.g., where the principal

⁶The set A_t is assumed to be one-dimensional for convenience; its openness is posited to exclude corner solutions for reasons explained in footnote 7.

⁷Openness of A_t is assumed to ensure that for all $\hat{\theta}_t$, θ_t , and a_t , there exists $\hat{a}_t$ such that $f_t(\theta_t, a_t) = f_t(\hat{\theta}_t, \hat{a}_t)$. If A_t were compact then there would be a pair, (θ'_t, a'_t) , that maximizes f_t . Therefore, if $\hat{\theta}_t \notin \arg \max_{\theta_t} [\max_{a_t} f_t(\theta_t, a_t)]$, then $f_t(\theta_t, a_t) = f_t(\hat{\theta}_t, \hat{a}_t)$ would not hold for any $\hat{a}_t$.

is a social planner and the agent the representative consumer) it could be the agent's expected payoff; in yet other applications it could be something different.

We denote partial derivatives with a subscript referring to the variable of differentiation, e.g., $\tilde{u}_{\theta_t} \equiv \partial \tilde{u} / \partial \theta_t$, $f_{i\theta_t} \equiv \partial f_i / \partial \theta_t$, etc.

Orthogonalization of information

The model can be transformed into an equivalent one where the agent's private information is represented by serially independent random variables. Suppose that at each $t = 0, \dots, T$, the agent observes $\varepsilon_t = G_t(\theta_t | \theta^{t-1}, a^{t-1}, x^{t-1})$ instead of θ_t . Clearly, ε^t can be inferred from $(\theta^t, a^{t-1}, x^{t-1})$. Conversely, θ^t can be computed from $(\varepsilon^t, a^{t-1}, x^{t-1})$, that is, for all $t = 0, \dots, T$, there is $\psi_t : [0, 1]^t \times A^{t-1} \times X^{t-1} \rightarrow \Theta_t$ such that

$$\varepsilon_t = G_t(\psi_t(\varepsilon^t, a^{t-1}, x^{t-1}) | \psi^{t-1}(\varepsilon^{t-1}, a^{t-2}, x^{t-2}), a^{t-1}, x^{t-1}),$$

where $\psi^t(\varepsilon^t, a^{t-1}, x^{t-1})$ denotes $(\psi_0(\varepsilon_0), \dots, \psi_t(\varepsilon^t, a^{t-1}, x^{t-1}))$. In other words, if the agent observes $(\varepsilon^t, a^{t-1}, x^{t-1})$ at time t in the orthogonalized model, she can infer the type history $\psi^t(\varepsilon^t, a^{t-1}, x^{t-1})$ in the original model.

Of course, a model where the agent observes ε_t for all t is strategically equivalent to the one where she observes θ_t for all t (provided that in both cases she observes x^{t-1} and recalls a^{t-1} at t). By definition, ε_t is uniformly distributed on the unit interval for all t and all realizations of θ^{t-1} , a^{t-1} , and x^{t-1} ; hence the random variables $\{\varepsilon_t\}_0^T$ are independent across time.⁸ There are many other orthogonalized information structures (e.g., those obtained by strictly monotonic transformations). In what follows, to simplify notation, we fix the orthogonalized information structure as that where ε_t is uniform on $\mathcal{E}_t = [0, 1]$.

The agent's gross payoff (i.e., utility before payments are subtracted) in the orthogonalized model, $u : \mathcal{E}^T \times A^T \times S^T \times X^T \rightarrow \mathbb{R}$, becomes

$$u(\varepsilon^T, a^T, s^T, x^T) = \tilde{u}(\psi^T(\varepsilon^T, a^{T-1}, x^{T-1}), a^T, s^T, x^T).$$

Revelation principle

A deterministic mechanism is a four-tuple $(Z^T, \mathbf{x}^T, \mathbf{a}^T, p)$, where Z_t is the agent's message space at time t , $\mathbf{x}_t : Z^t \times S^t \rightarrow X_t$ is the contractible decision rule at time t , $\mathbf{a}_t : Z^t \times S^{t-1} \rightarrow A_t$ is a recommended action at t , and $\mathbf{p} : Z^T \times S^T \rightarrow \mathbb{R}$ is the payment rule. The agent's reporting strategy at t is a mapping from previous reports and information to a message.

We refer to a strategy that maximizes the agent's payoff as an *equilibrium strategy* and the payoff generated by such a strategy as *equilibrium payoff*. The standard revelation principle applies in this setting, so it is without loss of generality to assume that

⁸To see that each ε_t is uniform on $[0, 1]$ conditional on the history of types, actions, and decisions up to t , note that since $\varepsilon_t = G_t(\theta_t | \theta^{t-1}, a^{t-1}, x^{t-1})$, the probability of $\varepsilon_t \leq \bar{\varepsilon}$ is $\Pr(G_t(\theta_t | \theta^{t-1}, a^{t-1}, x^{t-1}) \leq \bar{\varepsilon}) = \Pr(\theta_t \leq G_t^{-1}(\bar{\varepsilon} | \theta^{t-1}, a^{t-1}, x^{t-1})) = G_t(G_t^{-1}(\bar{\varepsilon} | \theta^{t-1}, a^{t-1}, x^{t-1}) | \theta^{t-1}, a^{t-1}, x^{t-1}) = \bar{\varepsilon}$.

$Z_t = \mathcal{E}_t$ for all t , and to restrict attention to mechanisms where telling the truth and taking the recommended action (obedience) is an equilibrium strategy. A direct mechanism is defined by a triple $(\mathbf{x}^T, \mathbf{a}^T, \mathbf{p})$, where $\mathbf{x}_t : \mathcal{E}^t \times S^t \rightarrow X_t$, $\mathbf{a}_t : \mathcal{E}^t \times S^{t-1} \rightarrow A_t$, and $\mathbf{p} : \mathcal{E}^T \times S^T \rightarrow \mathbb{R}$. Direct mechanisms in which telling the truth and obeying the principal's recommendation is an equilibrium strategy are called *incentive compatible* mechanisms.

We call a decision–action rule $(\mathbf{x}^T, \mathbf{a}^T)$ *implementable* if there exists a payment rule, $\mathbf{p} : \mathcal{E}^T \rightarrow \mathbb{R}$ such that the direct mechanism $(\mathbf{x}^T, \mathbf{a}^T, \mathbf{p})$ is incentive compatible.

Technical assumptions

We make three technical assumptions to ensure that the equilibrium payoff function of the agent is Lipschitz-continuous in the orthogonalized model.

ASSUMPTION 0. (i) *There exists a $K \in \mathbb{N}$ such that for all $t = 1, \dots, T$ and for all θ^T, a^T, s^T, x^T ,*

$$\tilde{u}_{\theta_t}(\theta^T, a^T, s^T, x^T), \tilde{u}_{a_t}(\theta^T, a^T, s^T, x^T) < K.$$

(ii) *There exists a $K \in \mathbb{N}$ such that for all $t = 1, \dots, T$, $\tau < t$, and for all $\theta^t, a^{t-1}, x^{t-1}$,*

$$G_{t\theta_t}(\theta_t | \theta^{t-1}, a^{t-1}, x^{t-1}), |G_{t\theta_\tau}(\theta_t | \theta^{t-1}, a^{t-1}, x^{t-1})| < K.$$

(iii) *There exists a $K \in \mathbb{N}$ such that for all $t = 1, \dots, T$ and for all θ^t, a^t , $f_{ta_t}(\theta_t, a_t) \neq 0$ and*

$$\left| \frac{f_{t\theta_t}(\theta_t, a_t)}{f_{ta_t}(\theta_t, a_t)} \right| < K.$$

3. THE MAIN RESULT

We refer to the model in which the principal never observes the agent's types as the *original model*, whereas we call the model where $\varepsilon_1, \dots, \varepsilon_T$ are observed by the principal the *benchmark case*. The contracting problem in the benchmark is static in the sense that the principal only needs the agent to report her information at $t = 0$. Our *irrelevance result* is that in any mechanism that implements a given decision–action rule in the original model, the principal pays the agent the same rent (i.e., he can achieve the same expected revenue) as in the benchmark case.

Specifically, what we show below is that the expected transfer payment of an agent with a given initial type when the principal implements decision–action rule $(\mathbf{x}^T, \mathbf{a}^T)$ in the original problem is the same (up to a type-invariant constant) as it would be in the benchmark. This implies that the principal can obtain the same expected revenue (or payoff, if it is linear in the expected revenue) when implementing a decision–action rule in the original problem as he could in the benchmark. This does not imply, however, that the two problems are equivalent: sufficient conditions of implementability (of a decision rule) are stronger in the original problem than they are in the benchmark. We will turn to the question of implementability in [Section 4](#).

In the next subsection we consider a decision–action rule $(\mathbf{x}^T, \mathbf{a}^T)$ and derive a necessary condition for the payment rule $\mathbf{p}$ such that $(\mathbf{x}^T, \mathbf{a}^T, \mathbf{p})$ is incentive compatible. This condition turns out to be the same in the benchmark case and in the original model. We then use this condition to prove our main, irrelevance result.

3.1 Payment rules

We fix an incentive compatible mechanism $(\mathbf{x}^T, \mathbf{a}^T, \mathbf{p})$ and analyze the consequences of time-0 incentive compatibility on the payment rule, $\mathbf{p}$, in both the original model and the benchmark.

We consider a particular set of deviation strategies and explore the consequence of the nonprofitability of these deviations in each case. To this end, let us define this set as follows: If the agent with initial type ε_0 reports $\widehat{\varepsilon}_0$, then (i) she must report $\varepsilon_1, \dots, \varepsilon_T$ truthfully and (ii) for all $t = 0, \dots, T$, after history (ε^t, s^{t-1}) , she must take action $\widehat{\mathbf{a}}_t(\varepsilon^t, \widehat{\varepsilon}_0, s^{t-1})$ such that the distribution of s_t is the same as if the history were $(\widehat{\varepsilon}_0, \varepsilon_{-0}^t, s^{t-1})$ and action $\mathbf{a}_t(\widehat{\varepsilon}_0, \varepsilon_{-0}^t, s^{t-1})$ were taken, where $\varepsilon_{-0}^t = (\varepsilon_1, \dots, \varepsilon_t)$. Since the distribution of s_t only depends on $f_t(\theta_t, \mathbf{a}_t)$, the action $\widehat{\mathbf{a}}_t(\varepsilon^t, \widehat{\varepsilon}_0, s^{t-1})$ is defined by

$$f_t(\widehat{\theta}_t, \mathbf{a}_t(\widehat{\varepsilon}_0, \varepsilon_{-0}^t, s^{t-1})) = f_t(\theta_t, \widehat{\mathbf{a}}_t(\varepsilon^t, \widehat{\varepsilon}_0, s^{t-1})), \quad (1)$$

where

$$\begin{aligned} \widehat{\theta}_t &= \psi_t(\widehat{\varepsilon}_0, \varepsilon_{-0}^t, \mathbf{a}^{t-1}(\widehat{\varepsilon}_0, \varepsilon_{-0}^{t-1}, s^{t-2}), \mathbf{x}^{t-1}(\widehat{\varepsilon}_0, \varepsilon_{-0}^{t-1}, s^{t-1})) \\ \theta_t &= \psi_t(\varepsilon^t, \widehat{\mathbf{a}}^{t-1}(\varepsilon^{t-1}, \widehat{\varepsilon}_0, s^{t-2}), \mathbf{x}^{t-1}(\widehat{\varepsilon}_0, \varepsilon_{-0}^{t-1}, s^{t-1})). \end{aligned}$$

In other words, the deviation strategies we consider require the agent (i) to be truthful in the future about her orthogonalized types and (ii) to take actions that “mask” her earlier lie so that the principal could not detect her initial deviation based on the contractible signals, even in a statistical sense.⁹ Note that in the benchmark case we only need to impose restriction (ii) since the principal observes $\varepsilon_1, \dots, \varepsilon_T$ by assumption. Also note that the strategies satisfying restrictions (i) and (ii) include the equilibrium strategy in the original model because if $\varepsilon_0 = \widehat{\varepsilon}_0$, the two restrictions imply truth-telling and obedience (adherence to the action rule).

We emphasize that we do not claim by any means that after reporting $\widehat{\varepsilon}_0$ it is *optimal* for the agent to follow a continuation strategy defined by restrictions (i) and (ii). Nevertheless, since the mechanism $(\mathbf{x}^T, \mathbf{a}^T, \mathbf{p})$ is incentive compatible, none of these deviations is profitable for the agent. We show that this observation enables us to characterize the expected payment of the agent conditional on ε_0 up to a type-invariant constant.

Let $\Pi_0(\varepsilon_0)$ denote the agent’s expected equilibrium payoff conditional on her initial type ε_0 in the incentive compatible mechanism $(\mathbf{x}^T, \mathbf{a}^T, \mathbf{p})$; that is,

$$\Pi_0(\varepsilon_0) = E[u(\varepsilon^T, \mathbf{a}^T(\varepsilon^T, s^{T-1}), s^T, \mathbf{x}^T(\varepsilon^T, s^T)) - \mathbf{p}(\varepsilon^T) | \varepsilon_0], \quad (2)$$

where E denotes expectation over ε^T and s^T .

⁹Similar ideas are used by Pavan et al. (2014) in a dynamic contracting model without hidden action and by Garrett and Pavan (2012) in a more restrictive environment with hidden action.

PROPOSITION 1. *If the mechanism $(\mathbf{x}^T, \mathbf{a}^T, \mathbf{p})$ is incentive compatible either in the original model or in the benchmark case, then, for all $\varepsilon_0 \in \mathcal{E}_0$,*

$$\begin{aligned} \Pi_0(\varepsilon_0) = \Pi_0(0) + E \left[\int_0^{\varepsilon_0} u_{\varepsilon_0}(y, \varepsilon_{-0}^T, \mathbf{a}^T(y, \varepsilon_{-0}^T, s^{T-1}), s^T, \mathbf{x}^T(y, \varepsilon_{-0}^T, s^T)) dy \middle| \varepsilon_0 \right] \\ + E \left[\int_0^{\varepsilon_0} \sum_{t=0}^T u_{a_t}(y, \varepsilon_{-0}^T, \mathbf{a}^T(y, \varepsilon_{-0}^T, s^{T-1}), s^T, \mathbf{x}^T(y, \varepsilon_{-0}^T, s^T)) \right. \\ \left. \times \widehat{\mathbf{a}}_{t\varepsilon_0}(y, \varepsilon_{-0}^t, y, s^{t-1}) dy \middle| \varepsilon_0 \right], \end{aligned} \quad (3)$$

where $(y, \varepsilon_{-0}^t) = (y, \varepsilon_1, \dots, \varepsilon_t)$.

Proposition 1 establishes that in an incentive compatible mechanism that implements a particular decision–action rule the expected payoff of the agent with a given (initial) type does not depend on the transfers. Analogous to the necessity part of the Spence–Mirrlees lemma in static mechanism design (or Myerson’s revenue equivalence theorem), necessary conditions similar to (3) have been derived in dynamic environments by Baron and Besanko (1984), Courty and Li (2000), Eső and Szentes (2007a), Pavan et al. (2014), Garrett and Pavan (2012), and others. In our environment, which is not only dynamic but incorporates both hidden information and hidden action as well, the real significance of the result is that *the same formula applies* in the original problem and in the benchmark case.¹⁰

It may be instructive to consider the special case where the principal has no access to contractible signals or, equivalently, the distribution of s_t is independent of (θ_t, a_t) . Since the choice of a_t has no impact on $\mathbf{x}^T$ and $\mathbf{p}^T$, the agent chooses a_t to maximize her utility. A necessary condition of this maximization is $E[u_{a_t}(\varepsilon^T, \mathbf{a}^T(\varepsilon^T, s^{T-1}), s^T, \mathbf{x}^T(\varepsilon^T, s^T)) | \varepsilon^t, s^{t-1}] = 0$ for all t . As a consequence the last term of $\Pi_0(\varepsilon_0)$, i.e., the second line of (3), vanishes.

PROOF OF PROPOSITION 1. First we express the agent’s reporting problem at $t = 0$ in the benchmark case as well as in the original problem subject to restrictions (i) and (ii) discussed at the beginning of this subsection.

To do this define

$$U(\varepsilon_0, \widehat{\varepsilon}_0) = E[u(\varepsilon^T, \widehat{\mathbf{a}}^T(\varepsilon^T, \widehat{\varepsilon}_0, s^{T-1}), s^T, \mathbf{x}^T(\widehat{\varepsilon}_0, \varepsilon_{-0}^T, s^T)) | \varepsilon_0]$$

and

$$P(\widehat{\varepsilon}_0) = E[\mathbf{p}(\widehat{\varepsilon}_0, \varepsilon_{-0}^T, s^T) | \varepsilon_0, a^T = \widehat{\mathbf{a}}^T(\varepsilon^T, \widehat{\varepsilon}_0, s^{T-1}), x^T = \mathbf{x}^T(\widehat{\varepsilon}_0, \varepsilon_{-0}^T, s^T)], \quad (4)$$

where $\widehat{\mathbf{a}}$ is defined by (1). Recall that the action $\widehat{\mathbf{a}}_t(\varepsilon^t, \widehat{\varepsilon}_0, s^{t-1})$ generates the same distribution of s_t as if the agent’s true type history was $(\widehat{\varepsilon}_0, \varepsilon_{-0}^t)$ and the agent had taken

¹⁰The derivation relies on the connectedness of the support of the type distributions. In a simpler environment, Krämer and Strausz (2015b) show that the irrelevance result fails with discrete types.

$\mathbf{a}_t(\widehat{\varepsilon}_0, \varepsilon_{-0}^t, s^{t-1})$. The significance of this is that

$$\begin{aligned} E[\mathbf{p}(\widehat{\varepsilon}_0, \varepsilon_{-0}^T, s^T) | \varepsilon_0, a^T = \widehat{\mathbf{a}}^T(\varepsilon^T, \widehat{\varepsilon}_0, s^{T-1}), x^T = \mathbf{x}^T(\widehat{\varepsilon}_0, \varepsilon_{-0}^T, s^T)] \\ = E[\mathbf{p}(\widehat{\varepsilon}_0, \varepsilon_{-0}^T, s^T) | \widehat{\varepsilon}_0, a^T = \mathbf{a}^T(\widehat{\varepsilon}_0, \varepsilon_{-0}^T, s^{T-1}), x^T = \mathbf{x}^T(\widehat{\varepsilon}_0, \varepsilon_{-0}^T, s^T)], \end{aligned}$$

so the right-hand side of (4) is indeed only a function of $\widehat{\varepsilon}_0$ but not that of ε_0 .

In the benchmark case, the payoff of the agent with ε_0 who reports $\widehat{\varepsilon}_0$ and takes action $\widehat{\mathbf{a}}_t(\varepsilon^t, \widehat{\varepsilon}_0, s^{t-1})$ at every t is $W(\varepsilon_0, \widehat{\varepsilon}_0) = U(\varepsilon_0, \widehat{\varepsilon}_0) - P(\widehat{\varepsilon}_0)$. Note that $W(\varepsilon_0, \widehat{\varepsilon}_0)$ is also the payoff of the agent in the original model if her type is ε_0 at $t = 0$, she reports $\widehat{\varepsilon}_0$, and her continuation strategy is defined by restrictions (i) and (ii) above, that is, she reports truthfully afterward and takes action $\widehat{\mathbf{a}}_t(\varepsilon^t, \widehat{\varepsilon}_0, s^{t-1})$ after the history (ε^t, s^{t-1}) .

Incentive compatibility of $(\mathbf{x}^T, \mathbf{a}^T, \mathbf{p})$ implies that $\varepsilon_0 \in \arg \max_{\widehat{\varepsilon}_0 \in \mathcal{E}_0} W(\varepsilon_0, \widehat{\varepsilon}_0)$ both in the benchmark case and in the original model. In addition, $\Pi_0(\varepsilon_0) = W(\varepsilon_0, \varepsilon_0)$ and, by Lemma 3 stated and proved in Appendix A, Π_0 is Lipschitz-continuous. Therefore, Theorem 1 in Milgrom and Segal (2002) implies that

$$\frac{d\Pi_0(\varepsilon_0)}{d\varepsilon_0} = \left. \frac{\partial U(\varepsilon_0, \widehat{\varepsilon}_0)}{\partial \varepsilon_0} \right|_{\widehat{\varepsilon}_0 = \varepsilon_0}$$

almost everywhere. (Differentiability of $\widehat{\mathbf{a}}_t$ in ε_0 is also established in the proof of Lemma 3.) Note that

$$\begin{aligned} & \left. \frac{\partial U(\varepsilon_0, \widehat{\varepsilon}_0)}{\partial \varepsilon_0} \right|_{\widehat{\varepsilon}_0 = \varepsilon_0} \\ &= E[u_{\varepsilon_0}(\varepsilon^T, \mathbf{a}^T(\varepsilon^T, s^{T-1}), s^T, \mathbf{x}^T(\varepsilon^T, s^T)) | \varepsilon_0] \\ & \quad + E \left[\sum_{t=0}^T u_{a_t}(\varepsilon^t, \mathbf{a}^t(\varepsilon^t, s^{t-1}), s^t, \mathbf{x}^t(\varepsilon^t, s^t)) \widehat{\mathbf{a}}_{t\varepsilon_0}(\varepsilon^t, \widehat{\varepsilon}_0, s^{t-1}) \right]_{\widehat{\varepsilon}_0 = \varepsilon_0} \Big|_{\varepsilon_0}. \end{aligned}$$

Since Π_0 is Lipschitz-continuous, it can be recovered from its derivative, so the statement of the proposition follows. $\square$

Note that the validity of the proof of Proposition 1 rests on the fact that the distribution of the agent's future, orthogonalized types $(\varepsilon_1, \dots, \varepsilon_T)$ does not depend on the realization of ε_0 ; otherwise there would be additional terms involving the derivatives (with respect to ε_0) of the conditional densities of future types in the expression for $\partial U(\varepsilon_0, \widehat{\varepsilon}_0) / \partial \varepsilon_0$.

By Proposition 1, for a given decision-action rule, incentive compatibility constraints pin down the expected payments conditional on ε_0 uniquely up to a constant in both the benchmark case and the original model. To see this, note that from (2) and (3) the expected payment conditional on ε_0 can be expressed as

$$\begin{aligned} & E[\mathbf{p}(\varepsilon^T, s^T) | \varepsilon_0, \mathbf{a}^T, \mathbf{x}^T] \\ &= E[u(\varepsilon^T, \mathbf{a}^T(\varepsilon^T, s^{T-1}), s^T, \mathbf{x}^T(\varepsilon^T)) | \varepsilon_0] - \Pi_0(0) \end{aligned}$$

$$\begin{aligned}
& -E \left[\int_0^{\varepsilon_0} u_{\varepsilon_0}(y, \varepsilon_{-0}^T, \mathbf{a}^T(y, \varepsilon_{-0}^T, s^{T-1}), s^T, \mathbf{x}^T(y, \varepsilon_{-0}^T, s^T)) dy \middle| \varepsilon_0 \right] \\
& -E \left[\int_0^{\varepsilon_0} \sum_{t=0}^T u_{a_t}(y, \varepsilon_{-0}^T, \mathbf{a}^T(y, \varepsilon_{-0}^T, s^{T-1}), s^T, \mathbf{x}^T(y, \varepsilon_{-0}^T, s^T)) \right. \\
& \quad \left. \times \widehat{\mathbf{a}}_{t\varepsilon_0}(y, \varepsilon_{-0}^t, y, s^{t-1}) \middle| \varepsilon_0 \right].
\end{aligned}$$

An immediate consequence of this observation and [Proposition 1](#) is the following remark.

REMARK 1. Suppose that $(\mathbf{x}^T, \mathbf{a}^T, \mathbf{p})$ and $(\mathbf{x}^T, \mathbf{a}^T, \bar{\mathbf{p}})$ are incentive compatible mechanisms in the original and in the benchmark case, respectively. Then there exists $c \in \mathbb{R}$ such that

$$E[\mathbf{p}(\varepsilon^T, s^T) | \varepsilon_0, \mathbf{a}^T, \mathbf{x}^T] - E[\bar{\mathbf{p}}(\varepsilon^T, s^T) | \varepsilon_0, \mathbf{a}^T, \mathbf{x}^T] = c.$$

3.2 Dynamic irrelevance

Now we show that the principal can achieve the same expected revenue implementing a decision rule *as if* he were able to observe the orthogonalized types of the agent after $t = 0$; that is, whenever a decision rule is implementable, the agent only receives information rents for her initial private information. This is our irrelevance result.

To state this result formally, suppose that the agent has an outside option, which we normalize to be zero. This means that any mechanism must satisfy

$$\Pi_0(\varepsilon_0) \geq 0, \quad \text{for all } \varepsilon_0 \in \mathcal{E}_0. \quad (5)$$

We call the maximum (supremum) of the expectation of payment $\mathbf{p}$ implementing $(\mathbf{x}^T, \mathbf{a}^T)$ and satisfying (5) the principal's *maximal revenue* from implementing $(\mathbf{x}^T, \mathbf{a}^T)$.¹¹

THEOREM 1. Suppose that decision rule $(\mathbf{x}^T, \mathbf{a}^T)$ is implementable in the original model. If payment rule $\bar{\mathbf{p}}$ implements $(\mathbf{x}^T, \mathbf{a}^T)$ in the benchmark case subject to (5), then there exists payment rule $\mathbf{p}$ that implements $(\mathbf{x}^T, \mathbf{a}^T)$ subject to (5) in the original model such that

$$E[\mathbf{p}(\varepsilon^T, s^T) | \varepsilon_0, \mathbf{a}^T, \mathbf{x}^T] = E[\bar{\mathbf{p}}(\varepsilon^T, s^T) | \varepsilon_0, \mathbf{a}^T, \mathbf{x}^T]. \quad (6)$$

In addition, the principal's maximal revenue from implementing $(\mathbf{x}^T, \mathbf{a}^T)$ in the original model is the same as in the benchmark case.

¹¹Requiring (5) for all ε_0 implies that we restrict attention to mechanisms where the agent participates irrespective of her type. This is without the loss of generality in many applications where there is a decision that generates a utility of zero for both the principal and the agent. Alternatively, we could have stated our theorem for problems where the participating types in the optimal contract of the benchmark case are an interval.

PROOF. Suppose that the direct mechanism $(\mathbf{x}^T, \mathbf{a}^T, \hat{\mathbf{p}})$ is incentive compatible in the original model. Then, by Remark 1,

$$E[\hat{\mathbf{p}}(\varepsilon^T, s^T)|\varepsilon_0, \mathbf{a}^T, \mathbf{x}^T] = E[\bar{\mathbf{p}}(\varepsilon^T, s^T)|\varepsilon_0, \mathbf{a}^T, \mathbf{x}^T] + c \quad (7)$$

for some $c \in \mathbb{R}$. Define $\mathbf{p}(\varepsilon^T, s^T)$ to be $\hat{\mathbf{p}}(\varepsilon^T, s^T) - c$. Since adding a constant has no effect on incentives, the mechanism $(\mathbf{x}^T, \mathbf{a}^T, \mathbf{p})$ is incentive compatible. In addition, the participation constraint of the agent, (5), is also satisfied because the agent's expected payoff conditional on her initial type, ε_0 , is the same as that in the benchmark case. Finally, notice that (7) implies (6), that is, the principal's revenue is the same as in the benchmark case.

It remains to argue that the principal's maximal revenue from implementing $(\mathbf{x}^T, \mathbf{a}^T)$ in the original model is the same as in the benchmark case. Note that the principal's expected revenue in the benchmark case is an upper bound on the same in the original model. We have just shown that if $\bar{\mathbf{p}}$ implements $(\mathbf{x}^T, \mathbf{a}^T)$ in the benchmark case, the principal can achieve $E[\bar{\mathbf{p}}(\varepsilon^T, s^T)|\mathbf{a}^T, \mathbf{x}^T]$ when implementing $(\mathbf{x}^T, \mathbf{a}^T)$ even if he does not observe $\varepsilon_1, \dots, \varepsilon_T$. $\square$

The statement of Theorem 1 is about the revenue of the principal. Note that if the payoff of the principal is also quasilinear (affine in the payment), then the decision rule and the expected payment fully determine his payoff. Hence, a consequence of Theorem 1 is the following remark.

REMARK 2. Suppose that the decision–action rule $(\mathbf{x}^T, \mathbf{a}^T)$ is implementable in the original model and the principal's payoff is affine in the payment. Then the principal's maximum (supremum) payoff from implementing $(\mathbf{x}^T, \mathbf{a}^T)$ is the same as in the benchmark case.

It is important to point out that our dynamic irrelevance result does not imply that the original problem (unobservable $\varepsilon_1, \dots, \varepsilon_T$) and the benchmark case (observable $\varepsilon_1, \dots, \varepsilon_T$) are equivalent. Theorem 1 only states that *if* an decision–action rule is implementable in the original model, *then* it can be done so without revenue loss as compared to the benchmark case. This result was obtained under very mild conditions regarding the stochastic process governing the agent's type, her payoff function, and the structure of signals. The obvious, next question is what type of decision–action rules can be implemented (under what conditions) in the original problem. In the next section we show that monotonicity of the decision–action rule is sufficient for implementation in certain environments. This result is used to solve applications in Section 5.

4. IMPLEMENTATION

This section establishes results regarding the implementability of certain decision rules.¹² We restrict attention to a Markov environment with time-separable, regular

¹²Throughout this section we require a type-invariant participation constraint for the agent with her outside option normalized to zero payoff, that is, we require (5) to hold.

(monotonic and single-crossing) payoff functions, formally stated in Assumptions 1 and 2 below.

First, we show that in the pure adverse selection model (where there are neither unobservable actions nor contractible signals) any monotonic decision rule is implementable. Then we turn our attention to the general model with moral hazard. There the set of implementable decision rules depends on the information content of the contractible signal. If the contractible signal has no informational content, that is, the distribution of s_t is independent of $f_t(\theta_t, a_t)$, then naturally the agent cannot be given incentives to choose any action other than the one that maximizes her flow utility in each period. In this case, we show that any decision–action rule can be implemented if $\mathbf{x}^T$ is monotonic and $\mathbf{a}^T$ is determined by the agent's per-period maximization problem.

The most interesting (and permissive) implementation result is obtained in the general model with adverse selection and moral hazard in case the signal is informative and its distribution satisfies a genericity condition due to McAfee and Reny (1992). This condition requires that the distribution of s_t conditional on any given $y_t = f_t(\theta_t, a_t)$ is not the average of signal distributions conditional on other $\hat{y}_t \neq y_t$'s such that $\hat{y}_t = f_t(\hat{\theta}_t, \hat{a}_t)$. In this case, we show that *any* monotonic decision–action rule $(\mathbf{x}^T, \mathbf{a}^T)$ can be *approximately* implemented (to be formally defined below).¹³ The result is based on arguments similar to the full surplus extraction theorem of McAfee and Reny (1992) and exploits the property of the model that f_t is approximately contractible and the agent is risk neutral with respect to monetary transfers. The main result of this section is that in our general model, in a regular Markovian environment with transferable utility and generic signals, the principal is able to implement any monotonic decision–action rule while not incurring any agency cost apart from the information rent due to the agent's initial private information.

So as to state the regularity assumptions made throughout the section, we return to the model without orthogonalization. Throughout this section, we assume that the contractible signal does not affect the agent's payoff directly and we remove s^T from the arguments of $\tilde{u}$, that is, $\tilde{u} : \Theta^T \times A^T \times X^T \rightarrow \mathbb{R}$. We make two sets of assumptions regarding the environment. The first set concerns the type distribution; the second set concerns the agent's payoff function.

ASSUMPTION 1 (Type distribution). (i) For all $t \in \{0, \dots, T\}$, the random variable θ_t is distributed according to a continuous c.d.f. $G_t(\cdot | \theta_{t-1})$ supported on an interval $\Theta_t = [\underline{\theta}_t, \bar{\theta}_t]$.

(ii) For all $t \in \{1, \dots, T\}$, $G_t(\cdot | \theta_{t-1}) \geq G_t(\cdot | \hat{\theta}_{t-1})$ whenever $\theta_{t-1} \leq \hat{\theta}_{t-1}$.

Part (i) of Assumption 1 states that the agent's type follows a Markov process, that is, the type distribution at time t only depends on the type at $(t - 1)$ and not on prior actions or decisions. In addition, the support of θ_t only depends on t , so any type on Θ_t can be realized irrespective of θ_{t-1} . Part (ii) states that the type distributions at time t

¹³The approximation can be dispensed with if the contractible signal is the summary statistic $f_t(\theta_t, a_t)$ itself.

are ordered according to first-order stochastic dominance. The larger is the agent's type at time $t - 1$, the more likely it is to be large at time t .

ASSUMPTION 2 (Payoff function). (i) *There exist $\{\tilde{u}_t\}_{t=0}^T$, $\tilde{u}_t : \Theta_t \times A_t \times X^t \rightarrow \mathbb{R}$ continuously differentiable, such that*

$$\tilde{u}(\theta^T, a^T, x^T) = \sum_{t=0}^T \tilde{u}_t(\theta_t, a_t, x^t).$$

(ii) *For all $t \in \{0, \dots, T\}$, $\tilde{u}_t$ is strictly increasing in θ_t .*

(iii) *For all $t \in \{0, \dots, T\}$, $\theta_t \in \Theta_t$, $a_t \in A_t$: $\tilde{u}_{t\theta_t}(\theta_t, a_t, x^t) \geq \tilde{u}_{t\theta_t}(\theta_t, a_t, \hat{x}^t)$ whenever $x^t \geq \hat{x}^t$.*

Part (i) of [Assumption 2](#) says that the agent's utility is additively separable over time, such that her flow utility at time t only depends on θ_t and a_t (and not on any prior information and action) besides all decisions taken at or before t . Part (ii) requires the flow utility to be monotonic in the agent's type. Part (iii) is the standard single-crossing property for the agent's type and the contractible decision.

We refer to the model as the one with *pure adverse selection* if $\tilde{u}_{ta_t} \equiv 0$ for all t and the distribution of s_t is independent of f_t . Next we state our implementation result for this case ([Proposition 2](#)). Then in [Sections 4.1](#) and [4.2](#) we return to the general model with moral hazard. In both scenarios regarding the informational content of signal s_t discussed above we reduce the problem of implementation to that in an appropriately defined pure adverse selection problem.

PROPOSITION 2. *Suppose that Assumptions 0, 1, and 2 hold in a pure adverse selection model. Then a decision rule, $\tilde{\mathbf{x}}^T, \tilde{\mathbf{x}}_t : \Theta^t \rightarrow X_t$, is implementable if $\tilde{\mathbf{x}}_t$ is increasing for all t .*

By [Corollary 2](#) of [Pavan et al. \(2014\)](#), [Assumptions 1](#) and [2](#) imply their integral monotonicity condition; slight differences between their and our technical assumptions notwithstanding, our [Proposition 2](#) appears to be an implication of their [Theorem 2](#). For completeness, a proof using techniques of [Eső and Szentes \(2007a\)](#) is provided in [Appendix B](#).¹⁴

4.1 Uninformative signals

Suppose that the contractible signal is uninformative (i.e., s_t is independent of f_t). We maintain the assumption that the payoff function of the agent is time-separable and satisfies [Assumption 2](#), but now the flow utility at time t is allowed to vary with a_t .

Recall that the action space of the agent at time t , A_t , was assumed to be an open interval of $\mathbb{R}$ in [Section 2](#). This assumption ensured that the agent could mask earlier lies (about her type) with her later hidden actions without being given away by signal s_t . Since there is no role for signal s_t in the case considered here, we can relax the require-

¹⁴At the end of this proof we also show that the principal can implement more allocations in the benchmark case than in the original model.

ment that A_t is open. In fact, so as to discuss the implementability of allocation rules that may involve boundary actions, we assume that $A_t = [\underline{a}_t, \bar{a}_t]$ is a compact interval throughout this subsection.

ASSUMPTION 3. For all $t \in \{0, \dots, T\}$, for all $\theta_t \in \Theta_t$, $a_t, \hat{a}_t \in A_t$, and $x^t, \hat{x}^t \in X^t$,

- (i) $\tilde{u}_{ta_t^2}(\theta_t, a_t, x^t) \leq 0$
- (ii) $\tilde{u}_{t\theta_t}(\theta_t, a_t, x^t) \geq \tilde{u}_{t\theta_t}(\theta_t, \hat{a}_t, \hat{x}^t)$ whenever $a_t \geq \hat{a}_t$
- (iii) $\tilde{u}_{ta_t}(\theta_t, a_t, x^t) \geq \tilde{u}_{ta_t}(\theta_t, a_t, \hat{x}^t)$ whenever $x^t \geq \hat{x}^t$.

Part (i) of the assumption states that the agent's payoff is a concave function of her action. This is satisfied in applications where the action of the agent is interpreted as an effort, and the cost of exerting effort is a convex function of the effort. Part (ii) states that the single-crossing assumption is also satisfied for the action. In the previous application, this means that the marginal cost of effort is decreasing in the agent's type. Part (iii) requires the single-crossing property to hold with respect to actions and decisions.

In what follows, we turn the problem of implementation in this environment with adverse selection and moral hazard into one of pure adverse selection. Since there is no contractible information about the agent's action, her action maximizes her payoff in each period and after each history; that is, if the agent has type θ_t and the history of decisions is x^t , then she takes an action that maximizes $\tilde{u}_t(\theta_t, a_t, x^t)$. Motivated by this observation, let us define the agent's new flow utility function at time t , $v_t : \Theta_t \times X^t \rightarrow \mathbb{R}$, to be

$$v_t(\theta_t, x^t) = \max_{a_t} \tilde{u}_t(\theta_t, a_t, x^t).$$

We will apply our implementation result for the pure adverse selection case ([Proposition 2](#)) to the setting where the flow utilities of the agent are $\{v_t\}_{t=0}^T$ while keeping in mind that the action of the agent in each period t maximizes $\tilde{u}_t$.

To this end, let $\bar{\mathbf{a}}_t(\theta_t, x^t)$ denote the generically unique $\arg \max_{a_t} \tilde{u}_t(\theta_t, a_t, x^t)$ for all $\theta_t \in \Theta_t$ and $x^t \in X^t$. By part (i) of [Assumption 3](#), if $\bar{\mathbf{a}}_t(\theta_t, x^t)$ is interior, it is defined by the first-order condition

$$\tilde{u}_{ta_t}(\theta_t, \bar{\mathbf{a}}_t(\theta_t, x^t), x^t) = 0. \quad (8)$$

The next lemma states that the flow utilities, $\{v_t\}_0^T$, satisfy the hypothesis of [Proposition 2](#).

LEMMA 1. Suppose that the functions $\{\tilde{u}_t\}_{t=0}^T$ satisfy [Assumptions 2 and 3](#). Then the functions $\{v_t\}_{t=0}^T$ satisfy [Assumption 2](#).

Suppose that the decision–action rule $(\mathbf{x}^T, \mathbf{a}^T)$ is implementable. Then, since the agent's action maximizes her payoff in each period, $\mathbf{a}_t(\theta^t) = \bar{\mathbf{a}}_t(\theta_t, \mathbf{x}^t(\theta^t))$. In addition, the decision rule $\mathbf{x}^T$ must be implementable in the pure adverse selection model, where the agent's flow utility functions are $\{v_t\}_{t=1}^T$. Hence, the following result is a consequence of [Proposition 2](#) and [Lemma 1](#).

PROPOSITION 3. *Suppose that Assumptions 0–3 hold. Then a decision rule, $(\tilde{\mathbf{x}}^T, \tilde{\mathbf{a}}^T)$, $\tilde{\mathbf{x}}_t : \Theta^t \rightarrow X_t$ and $\tilde{\mathbf{a}}_t : \Theta^t \rightarrow A_t$, is implementable if $\tilde{\mathbf{x}}_t$ is increasing and $\tilde{\mathbf{a}}_t(\theta^t) = \bar{\mathbf{a}}_t(\theta_t, \tilde{\mathbf{x}}^t(\theta^t))$ for all $t \in \{0, \dots, T\}$.*

Of course, the statement of this proposition is valid even if the contractible signal is informative (s_t depends on f_t) but the principal ignores it and designs a mechanism that does not condition on s^T . However, if s_t is informative about $f_t(\theta_t, a_t)$ the principal can implement more decision rules, which is the subject of the next subsection.

4.2 Informative signals

We turn our attention to the case where the contractible signal is informative. The next condition is due to McAfee and Reny (1992); it requires that the distribution of the contractible signal conditional of any given value of $y_0 = f_t(\theta_t, a_t)$ is *not* the average of the distribution of s_t conditional on other values of f_t . This condition is generic.

ASSUMPTION 4. *Suppose that for all $\theta_t \in \Theta_t$ and $a_t \in A_t$, $f_t(\theta_t, a_t) \in Y_t = [\underline{y}_t, \bar{y}_t]$. Then, for all $\mu \in \Delta[\underline{y}_t, \bar{y}_t]$ and $y_0 \in [\underline{y}_t, \bar{y}_t]$, $\mu(\{y_0\}) \neq 1$ implies $h(\cdot|y_0) \neq \int_{\underline{y}_t}^{\bar{y}_t} h(\cdot|y)\mu(dy)$.*

Next, we make further assumptions on the agent's flow utility, $\tilde{u}_t$, and on the shape of the function f_t .

ASSUMPTION 5. *For all $t \in \{0, \dots, T\}$, for all $\theta_t \in \Theta_t$, $a_t \in A_t$, $x^t \in X^t$,*

- (i) $\tilde{u}_{ta_t}(\theta_t, a_t, x^t) < 0$
- (ii) *there exists a $K \in \mathbb{N}$ such that $f_{ta_t}(\theta_t, a_t), f_{t\theta_t}(\theta_t, a_t) > 1/K$*
- (iii) $f_{ta_t^2}(\theta_t, a_t)f_{t\theta_t}(\theta_t, a_t) \leq f_{ta_t}(\theta_t, a_t)f_{ta_t\theta_t}(\theta_t, a_t)$
- (iv) $\tilde{u}_{t\theta_t x^t}(\theta_t, a_t, x^t)f_{ta_t}(\theta_t, a_t) \geq \tilde{u}_{ta_t x^t}(\theta_t, a_t, x^t)f_{t\theta_t}(\theta_t, a_t)$.

Part (i) requires the agent's flow utility to be decreasing in her action. This is satisfied in applications where, for example, the agent's unobservable action is a costly effort from which she does not benefit directly. Part (ii) says that the function f_t is increasing in both the agent's action and type. In many applications, the distribution of the contractible signal can be ordered according to first-order stochastic dominance. In these applications, part (ii) implies that an increase in either the action or the type improves the distribution of s_t in the sense of first-order stochastic dominance. Part (iii) is a substitution assumption regarding the agent's type and hidden action in the value of f_t . It means that an increase in a_t , holding the value of f_t constant, weakly decreases the *marginal* impact of a_t on f_t .¹⁵ This assumption is satisfied, for example, if $f_t(\theta_t, a_t) = \theta_t + a_t$, but it is clearly more general. As will be explained later, part (iv) is

¹⁵To see this interpretation, note that the total differential of f_{ta_t} (the change in the marginal impact of a_t) is $f_{ta_t^2} da_t + f_{ta_t\theta_t} d\theta_t$. Keeping f_t constant (moving along an "isovalue" curve) means $d\theta_t = (-f_{ta_t}/f_{t\theta_t}) da_t$. Substituting this into the total differential of f_{ta_t} yields $(f_{ta_t^2} - f_{ta_t\theta_t}f_{ta_t}/f_{t\theta_t}) da_t$. This expression is nonpositive for $da_t > 0$ if part (iii) is satisfied.

a strengthening of the single-crossing property posited in part (iii) of [Assumption 2](#). It requires the marginal utility in type to be increasing in the contractible decision while holding the value of f_t fixed. This assumption is satisfied, for example, if the effort cost of the agent is additively separable in her flow utility.

The key observation is that due to [Assumption 4](#), the value of f_t becomes an approximately contractible object in the following sense. For each value of f_t , y_t , the principal can design a transfer scheme depending only on s^T that punishes the agent for taking an action that results in a value of f_t that is different from y_t . Perhaps more importantly, the punishment can be arbitrarily large as a function of the distance between y_t and the realized value of f_t . We use this observation to establish our implementation result in two steps. First, we treat f_t (for all t) as a contractible object, that is, we add another dimension to the contractible decisions in each period. Since, conditional on θ_t , the value of f_t is determined by a_t , we can express the agent's flow utility as a function of f_t instead of a_t . These new flow utilities depend only on types and decisions, so we have a pure adverse selection model. We then show that the new flow utilities satisfy the requirements of [Proposition 2](#) and hence, every monotonic rule is implementable. The second step is to construct the punishment transfers mentioned above and show that even if f_t is not contractible, any monotonic decision rule can be approximately implementable.

For each $y_t \in \{f_t(\theta_t, a_t) : \theta_t \in \Theta_t, a_t \in A_t\}$ and $\theta_t \in \Theta_t$, let $\mathbf{a}_t(\theta_t, y_t)$ denote the solution to $f_t(\theta_t, a_t) = y_t$ in a_t . For each $t = 0, \dots, T$, we define the agent's flow utility as a function of y_t as

$$w_t(\theta_t, y_t, x^t) = \tilde{u}_t(\theta_t, \mathbf{a}_t(\theta_t, y_t), x^t).$$

Next, we show that the functions $\{w_t\}_{t=0}^T$ satisfy the hypothesis of [Proposition 2](#).

LEMMA 2. *Suppose that Assumptions 2–5 are satisfied. Then the functions $\{w_t\}_{t=0}^T$ satisfy [Assumption 2](#).*

By this lemma and [Proposition 2](#), if the value of f_t was contractible for all t , any increasing decision rule was implementable. However, f_t is not contractible; nevertheless we can still implement increasing decisions rules *approximately* in the sense that by following the principal's recommendation the agent's expected utility is arbitrarily close to her equilibrium payoff. The following definition gives this concept formally.

DEFINITION 1. The decision rule $(\tilde{\mathbf{x}}^T, \tilde{\mathbf{a}}^T)$ is approximately implementable if for all $\delta > 0$ there exists a payment rule $\tilde{\mathbf{p}} : \Theta^T \times S^T \rightarrow \mathbb{R}$ such that for all $\theta_0 \in \Theta_0$,

$$E_{s^T} \left[\sum_{t=0}^T \tilde{u}_t(\theta_t, \tilde{\mathbf{a}}_t(\theta^t), \tilde{\mathbf{x}}^t(\theta^t)) - \tilde{\mathbf{p}}(\theta^T, s^T) \middle| \theta_0 \right] \geq \Pi_0(\theta_0) - \delta, \quad (9)$$

where $\Pi_0(\theta_0)$ denotes the agent's equilibrium payoff with initial type θ_0 .

We are ready to state the implementation result of this subsection.

PROPOSITION 4. *Suppose that Assumptions 0–5 are satisfied. Then a decision rule, $(\tilde{\mathbf{x}}^T, \tilde{\mathbf{a}}^T)$, $\tilde{\mathbf{x}}_t : \Theta^t \rightarrow X_t$ and $\tilde{\mathbf{a}}_t : \Theta^t \rightarrow A_t$, is approximately implementable if $\tilde{\mathbf{x}}_t$ and $\tilde{\mathbf{a}}_t$ are increasing for all $t \in \{0, \dots, T\}$.*

In the proof of this proposition, we decomposed the gain from any deviation strategy into two parts. The first part is the difference between the payoff from truth-telling and deviating in the hypothetical model where y_t is contractible. The second part is the difference between the payoff from the misreporting strategy and taking actions corresponding to the misreports and the payoff from misreporting and taking actions optimally. Then we use [Proposition 2](#) to construct payments so that the first part is negative and use the payment construction of [McAfee and Reny \(1992\)](#) to guarantee that the second part is small. If we were able to show that the loss due to a misreporting in the hypothetical model where y_t is contractible is large if the deviation results in a decision rule that is far away from the intended decision rule, then we can prove that any profitable deviation results in a decision rule that is nearby the intended one. Since the optimal strategy in any mechanism is implementable, it would imply that there is an implementable decision rule nearby any increasing decision rule. So as to get a bound on deviation payoffs, just like in static mechanism design, we need to require the payoff function w_t to satisfy the strict single-crossing property. It turns out that the strict single-crossing property is satisfied if part (iii) [Assumption 2](#) and part (iv) of [Assumption 5](#) hold with strict inequalities.

PROPOSITION 5. *Suppose that Assumptions 0–5 are satisfied and the inequalities of part (iii) of [Assumption 2](#) and part (iv) of [Assumption 5](#) are strict. If the decision rule, $(\tilde{\mathbf{x}}^T, \tilde{\mathbf{a}}^T)$, $\tilde{\mathbf{x}}_t : \Theta^t \rightarrow X_t$ and $\tilde{\mathbf{a}}_t : \Theta^t \rightarrow A_t$, is continuous and increasing, then for all $\delta > 0$ there is an allocation rule, $(\bar{\mathbf{x}}^T, \bar{\mathbf{a}}^T)$, such that $(\bar{\mathbf{x}}^T, \bar{\mathbf{a}}^T)$ is implementable and*

$$E_{\theta^T} \sum_{t=0}^T \left\| (\tilde{\mathbf{y}}_t(\theta^t), \tilde{\mathbf{x}}_t(\theta^t)) - (\bar{\mathbf{y}}_t(\theta^t), \bar{\mathbf{x}}_t(\theta^t)) \right\|^2 < \delta,$$

where $\tilde{\mathbf{y}}_t(\theta^t) = f_t(\theta_t, \tilde{\mathbf{a}}_t(\theta^t))$ and $\bar{\mathbf{y}}_t(\theta^t) = f_t(\theta_t, \bar{\mathbf{a}}_t(\theta^t))$.

See Appendix B for the proof.

The implementation results of this section allow us to use a simple (and familiar) method for solving the contracting problem of a principal whose payoff is linear in the expected transfer. This method will be further explained in [Section 5](#) in the context of applications; here we give a brief summary for contracting problems satisfying the conditions (Markovian types, regular, time-separable utilities, and imperfect s_t signals) of [Proposition 4](#). First, suppose that y_t (the summary statistic about the agent's period- t type and action) is contractible. [Theorem 1](#) applies in this case; hence the maximal transfer in any contract is the same as it would be in the benchmark (where the principal observes the agent's orthogonalized types for all $t > 0$). Solve the benchmark problem with felicity functions $\{w_t\}$, as if y were contractible. If the resulting decision–action rule is monotonic in the agent's type profile, then by [Proposition 4](#) the same can be implemented approximately (with approximate incentive compatibility) even when the principal does not observe $(\varepsilon_1, \dots, \varepsilon_T)$ and only observes an imperfect signal s_t about each y_t . Moreover, as the proof of [Proposition 4](#) shows, the expected transfer from the approximate implementation is still the same as it would be with contractible $\{y_t\}$. (This is so because the expected value of the additional transfers is zero.) Hence the solution to the

benchmark case with felicities $\{w_t\}$ and observable $\{y_t\}$ is indeed the optimal mechanism in the original problem, provided the optimal decision–action rule is monotonic.

5. APPLICATIONS

We present three applications to illustrate how our techniques and results can be applied in substantive economic problems. In each application we first solve the benchmark case, where the principal can observe the agent's orthogonalized future types. (In the absence of a contractible summary signal about the agent's type and hidden action, the action rule is taken to be the agent-optimal one; in the presence of such a signal the action rule is also optimized.) Then we verify the appropriate monotonicity condition regarding the decision–action rule and conclude that the solution is implementable, hence optimal, in the original problem as well.

In all three applications we assume that the agent's type follows the (autoregressive) AR(1) process

$$\theta_t = \lambda \theta_{t-1} + (1 - \lambda) \varepsilon_t \quad \forall t = 0, \dots, T,$$

where $\theta_{-1} = 0$ and $\varepsilon_0, \dots, \varepsilon_T$ are independent and identically distributed (iid) uniform on $[0, 1]$. The exact specification is adopted for the sake of obtaining a simple orthogonal transformation of the information structure:

$$\theta_t = (1 - \lambda) \lambda^t \sum_{k=0}^t \lambda^{-k} \varepsilon_k \quad \forall t = 0, \dots, T. \quad (10)$$

The type process is Markovian. [Assumption 1](#) is satisfied except that the support of θ_t depends on the realization of θ_{t-1} . However, it is easy to make the support of θ_t the unit interval for all t by mixing the distribution of θ_t in (10) with the uniform distribution on $[0, 1]$; our specification obtains in the limit as the weight on the uniform distribution vanishes.

In all three examples the agent's utility is time-separable, and the flow utility, $\tilde{u}_t(\theta_t, a_t, x_t)$, only depends on the agent's type, hidden action, and the contractible decision.¹⁶ Denote the flow utility in the orthogonally transformed model by $u_t(\varepsilon^t, a_t, x_t)$. By [Proposition 1](#), in any incentive compatible mechanism $(\mathbf{x}^T, \mathbf{a}^T, p)$ the agent's equilibrium payoff can be written as

$$\begin{aligned} \Pi_0(\varepsilon_0) = \Pi_0(0) + E \left[\int_0^{\varepsilon_0} \sum_{t=0}^T u_{t\varepsilon_0}(y, \varepsilon_{-0}^t, a_t(y, \varepsilon_{-0}^t, s^{t-1}), x_t(y, \varepsilon_{-0}^t, s^t)) dy \middle| \varepsilon_0 \right] \\ + E \left[\int_0^{\varepsilon_0} \sum_{t=0}^T u_{ta_t}(y, \varepsilon_{-0}^t, a_t(y, \varepsilon_{-0}^t, s^{t-1}), x_t(y, \varepsilon_{-0}^t, s^t)) \right. \\ \left. \times \widehat{a}_{t\varepsilon_0}(y, \varepsilon_{-0}^t, y, s^{t-1}) dy \middle| \varepsilon_0 \right], \end{aligned} \quad (11)$$

¹⁶[Assumption 0](#) is also satisfied due to the boundedness of all relevant domains and the continuous differentiability of all involved functions.

where $\widehat{a}_t(\varepsilon^t, \widehat{\varepsilon}_0, s^{t-1})$, defined by (1), is the period- t action of the agent that “masks” her initial misreport of $\widehat{\varepsilon}_0$ conditional on the history of types and contractible signals.

Next, we describe Applications 1–3, ordered according to increasing complexity of the agent’s payoff function. The first application is a pure adverse selection model; the second one is a variant that includes a hidden action as well, but no contractible signal about the agent’s type and action. The third application has both hidden type and hidden action; the agent’s type and action generate a noisy but contractible summary signal.

APPLICATION 1. The principal is the seller of an indivisible good; the agent is a buyer with valuation θ_t in period t . The contractible action, $x_t \in [0, 1]$, is the probability that the buyer receives the good. The buyer has no hidden action; her flow utility is simply $\tilde{u}_t(\theta_t, x_t) = \theta_t x_t$ or, equivalently in the orthogonalized model, $u_t(\varepsilon^t, x_t) = (1 - \lambda)\lambda^t(\sum_{k=0}^t \lambda^{-k} \varepsilon_k)x_t$. Note that [Assumption 2](#) holds and $u_{t\varepsilon_0} = (1 - \lambda)\lambda^t x_t$.

Since the agent has no hidden action the second line in (11) is zero, and so

$$\Pi_0(\varepsilon_0) = \Pi_0(0) + E \left[\int_0^{\varepsilon_0} \sum_{t=0}^T (1 - \lambda)\lambda^t x_t(y, \varepsilon_{-0}^t) dy \middle| \varepsilon_0 \right]. \quad (12)$$

Suppose the buyer’s participation is guaranteed if she gets a nonnegative payoff; by (12) this is equivalent to $\Pi_0(0) \geq 0$.

So as to compute $E[\Pi_0(\varepsilon_0)]$, we note that by Fubini’s theorem,

$$\int_0^1 \int_0^{\varepsilon_0} x_t(y, \varepsilon_{-0}^t) dy d\varepsilon_0 = \int_0^1 \int_y^1 x_t(y, \varepsilon_{-0}^t) d\varepsilon_0 dy = \int_0^1 (1 - \varepsilon_0)x_t(\varepsilon^t) d\varepsilon_0;$$

therefore

$$E[\Pi_0(\varepsilon_0)] = \Pi_0(0) + E \left[\sum_{t=0}^T (1 - \lambda)\lambda^t (1 - \varepsilon_0)x_t(\varepsilon^t) \right]. \quad (13)$$

Assume the seller (principal) maximizes his expected revenue; there is no cost of production. The expected revenue equals the expected social surplus generated by the mechanism less the buyer’s expected payoff,

$$\sum_{t=0}^T E[\theta_t x_t(\varepsilon^t) - (1 - \lambda)\lambda^t (1 - \varepsilon_0)x_t(\varepsilon^t)] - \Pi_0(0),$$

where θ_t is given by (10). Solve the seller’s problem by setting $\Pi_0(0) = 0$ and pointwise maximizing the objective in $x_t(\varepsilon^t)$: the solution is found by setting $x_t^*(\varepsilon^t) = 1$ if and only if $\theta_t \geq (1 - \lambda)\lambda^t (1 - \varepsilon_0)$ and $x_t^*(\varepsilon^t) = 0$ otherwise. Equivalently, in the notation of the original model,

$$\tilde{x}^*(\theta^t) = \mathbf{1}_{\theta_t + \lambda^t \theta_0 \geq (1 - \lambda)\lambda^t},$$

where $\mathbf{1}$ is the indicator function. This decision rule is monotone in θ^t ; therefore, by [Proposition 2](#), it is implementable in the original problem as well as in the benchmark case. Hence it is the optimal solution in both.

In this multiperiod trading (single-buyer auction) problem the first-best outcome would be to trade the good whenever $\theta_t \geq 0$. In contrast, in the revenue-maximizing mechanism the good is sold whenever $\theta_t \geq \lambda^t(1 - \lambda - \theta_0)$. As in the one-period problem, this decision rule corresponds to setting a reservation price in each period. The reservation price is always nonnegative because $\theta_0 \leq 1 - \lambda$ by (10). Interestingly, the reservation prices and the distortion that they induce only depend on the buyer's initial information (confirming our dynamic irrelevance result) and disappear over time as $t \rightarrow \infty$.

APPLICATION 2. In this application, as in the previous one, the principal is a seller and the agent is a buyer with period- t valuation θ_t . Assume the good is divisible, so $x_t \in [0, 1]$ is interpreted as the amount bought by the buyer, and the seller has production cost $x_t^2/2$.

The important difference in this application (as compared to the previous one) is that we assume the buyer takes a costly, hidden action interpreted as investment in every period, which increases her valuation.¹⁷ The buyer's flow utility is $\tilde{u}_t(\theta_t, a_t, x_t) = (\theta_t + a_t)x_t - ca_t^2/2$ or, equivalently in the orthogonalized model,

$$u_t(\varepsilon^t, a_t, x_t) = \left[(1 - \lambda)\lambda^t \sum_{k=0}^t \lambda^{-k} \varepsilon_k + a_t \right] x_t - \frac{1}{2}ca_t^2.$$

Note that Assumptions 0–3 hold, and $u_{t\varepsilon_0} = (1 - \lambda)\lambda^t x_t$ (same as in [Application 1](#)).

Assume that the seller cannot observe any signal about the buyer's valuation and investment. Hence the second line in (11) is zero, and so $\Pi_0(\varepsilon_0)$ is given by (12) and $E[\Pi_0(\varepsilon_0)]$ is given by (13). The seller's (principal's) expected profit is the expected social surplus generated by the mechanism less the buyer's (agent's) expected payoff:

$$\sum_{t=0}^T E \left[(\theta_t + a_t(\varepsilon^t))x_t(\varepsilon^t) - \frac{1}{2}ca_t(\varepsilon^t)^2 - \frac{1}{2}x_t(\varepsilon^t)^2 - (1 - \lambda)\lambda^t(1 - \varepsilon_0)x_t(\varepsilon^t) \right] - \Pi_0(0).$$

Since the seller can make no inference about a_t and, moreover, the buyer's future valuations are not affected by her current investment either, a_t is set by the buyer to maximize her current flow utility: $a_t(\varepsilon^t) \equiv x_t(\varepsilon^t)/c$. Substituting this into the seller's expected payoff, the first-order condition of pointwise maximization of the seller's objective in $x_t(\varepsilon^t)$ is

$$\theta_t + \frac{x_t(\varepsilon^t)}{c} - x_t(\varepsilon^t) - (1 - \lambda)\lambda^t(1 - \varepsilon_0) = 0. \quad (14)$$

Assuming that the buyer participates with nonnegative payoff, it is optimal to set $\Pi_0(0) = 0$. Using (10) in rearranging (14) yields, in terms of the original model,

$$\tilde{x}_t^*(\theta^t) = \frac{c}{c-1} [\theta_t + \lambda^t \theta_0 - (1 - \lambda)\lambda^t].$$

¹⁷Interpreting a_t as a costly action taken right before θ_t is realized and shifting the distribution of θ_t , this application can be thought of as a multiperiod generalization (of a specific example) of [Bergemann and Välimäki \(2002\)](#). Our focus is on the revenue-maximizing sales mechanism instead of the efficient one.

Assume $c > 1$. Then $\tilde{x}_t^*(\theta^t)$ is strictly increasing; by [Proposition 3](#) it is implementable both in the original problem and the benchmark when coupled with investments $\tilde{a}^*(\theta^t) = \tilde{x}_t^*(\theta^t)/c$. Therefore this allocation rule is the optimal second-best solution in both problems.

In this application, in the first-best case (contractible θ_t, a_t), the relationship between the buyer's investment level and her anticipated purchase (trade) would be the same, $a_t^{\text{FB}} = x_t^{\text{FB}}/c$. However, the first-best level of trade would be $x_t^{\text{FB}}(\theta_t) = c\theta_t/(c-1)$. The distortion, which materializes in the decision rule in the form of less trade and in the action rule as less investment in comparison to the efficient levels, is again due to the buyer's (agent's) initial private information and it disappears over time.

APPLICATION 3. The principal is an investor (a wealthy individual or institution) and the agent is an investment advisor (private banker); the contractible action x_t is the amount invested, according to the agent's advice, on behalf of the principal. The agent's type θ_t represents her ability to achieve a greater expected return. Her costly effort (hidden action a_t) is directed at finding assets that fit the principal's other (e.g., ethical) investment goals; it generates a payoff proportional to the invested amount for the principal but imposes an up-front cost on the agent.

Let $\tilde{u}_t = -ca_t^2/2$ be the agent's payoff and let $v_t = (\theta_t + a_t + \xi_t)x_t - rx_t^2/2$ be the principal's payoff; in the latter $rx_t^2/2$ represents the principal's (convex) cost of raising funds for investment, and ξ_t is a noise term (e.g., uncertainty in how the advisor's effort affects the investor's nonpecuniary return on investment). Assume that v_t (but not θ_t or a_t) is contractible, and define $s_t = \theta_t + a_t + \xi_t$ as the contractible signal. The parties' payoffs are transferable, i.e., they may contract on monetary transfers as well. It is easy to check that in this application Assumptions 0–5 are all satisfied.¹⁸ This is a parametric example of the model discussed in [Section 4.2](#). [Garrett and Pavan \(2012\)](#) solve a related problem where, using the notation of this example, we have $r = 0$ and the decision $x_t \in \{0, 1\}$ corresponds to whether the principal employs the agent instead of a continuous investment decision (which is more meaningful in our application).

In the orthogonalized model $u_t(\varepsilon^t, a_t, x_t) = -ca_t^2/2$; hence $u_{t\varepsilon_0} = 0$ and $u_{ta_t} = -ca_t$. The period- t action of the agent that masks her initial misreport of $\hat{\varepsilon}_0$ conditional on the history of types and signals is $\hat{a}_t(\varepsilon^t, \hat{\varepsilon}_0, s^{t-1})$, formally defined by

$$\hat{\theta}_t + a_t(\hat{\varepsilon}_0, \varepsilon_{-0}^t, s^{t-1}) \equiv \theta_t + \hat{a}_t(\varepsilon^t, \hat{\varepsilon}_0, s^{t-1}),$$

where $\hat{\theta}_t = (1-\lambda)\lambda^t\hat{\varepsilon}_0 + (1-\lambda)\sum_{k=1}^t \lambda^{t-k}\varepsilon_k$; hence

$$\hat{a}_t(\varepsilon^t, \hat{\varepsilon}_0, s^{t-1}) = a_t(\hat{\varepsilon}_0, \varepsilon_{-0}^t, s^{t-1}) + (1-\lambda)\lambda^t(\hat{\varepsilon}_0 - \varepsilon_0).$$

Note that $\hat{a}_{t\varepsilon_0}(\varepsilon^t, \hat{\varepsilon}_0, s^{t-1}) = -(1-\lambda)\lambda^t$.

By (11), the agent's expected payoff with initial type ε_0 is

$$\Pi_0(\varepsilon_0) = \Pi_0(0) + E \left[\int_0^{\varepsilon_0} \sum_{t=0}^T (1-\lambda)\lambda^t ca_t(y, \varepsilon_{-0}^t, s^{t-1}) dy \middle| \varepsilon_0 \right].$$

¹⁸ [Assumption 2\(ii\)](#) holds weakly, but the approximate implementation result holds. The model could easily be generalized in a way that the agent's flow utility directly, positively depended on θ_t .

Continue to use $\Pi_0(0) \geq 0$ as the participation constraint. Again, using Fubini's theorem as in the previous applications we get

$$E[\Pi_0(\varepsilon_0)] = \Pi_0(0) + E \left[\sum_{t=0}^T (1-\lambda)\lambda^t(1-\varepsilon_0)ca_t(\varepsilon^t, s^{t-1}) \right].$$

The principal's ex ante expected payoff is the difference between the expected social surplus generated by the mechanism and the agent's expected payoff,

$$\sum_{t=0}^T E \left[(\theta_t + a_t + \xi_t)x_t - \frac{1}{2}rx_t^2 - \frac{1}{2}ca_t^2 - (1-\lambda)\lambda^t(1-\varepsilon_0)ca_t \right] - \Pi_0(0), \quad (15)$$

where the arguments of $a_t(\theta^t, s^{t-1})$ and $x_t(\theta^t, s^t)$ are suppressed for brevity.

If signal s_t contained no noise term (i.e., in case $\xi_t \equiv 0$), then the principal could infer a_t from the agent's type report and the realized signal, and indirectly enforce any action. In this case, the first-order condition of (pointwise) maximization of (15) in a_t is $x_t - ca_t - (1-\lambda)\lambda^t(1-\varepsilon_0)c = 0$, whereas the same with respect to x_t is $\theta_t + a_t - rx_t = 0$. Combine the two equations and write $\theta_0/(1-\lambda)$ for ε_0 to get

$$\tilde{x}_t^*(\theta^t) = \frac{\theta_t + \lambda^t\theta_0 - (1-\lambda)\lambda^t}{rc - 1}.$$

Assuming $rc > 1$ the resulting $\tilde{x}_t^*$ is strictly increasing in θ^t , and hence so is the corresponding optimal $\tilde{a}_t^*$, which is its positive affine transformation. Therefore, by [Proposition 4](#), this decision–action rule is approximately implementable in the original model as well as in the benchmark. It is easy to see that in the first-best case, $x_t^{\text{FB}}(\theta^t) = \theta_t/(rc - 1)$. Again, the distortion in $\tilde{x}^*(\theta^t)$ is purely due to the agent's initial private information, illustrating our dynamic irrelevance theorem.

6. CONCLUSIONS

In this paper we considered a dynamic principal–agent model with adverse selection and moral hazard, and proved a dynamic irrelevance theorem: In any fixed, implementable decision–action rule the principal's expected revenue and the agent's payoff are the same *as if* the principal could observe the agent's future, orthogonalized types. This result comes with (at least) two caveats: (i) the set of rules that can be implemented with or without observing the agent's future, orthogonalized types is different; (ii) the result pins down the expected payments at the time of contracting, but not their distribution over time. We also provided results on the implementability of monotonic decision rules in regular, Markovian environments. The implementation results imply a straightforward method of solving a large class of dynamic principal–agent problems with meaningful economic applications.

The model considered in this paper could be extended in two directions without much difficulty, at the expense of additional notation and technical assumptions. First, it would be possible to accommodate multiple agents in the principal–agent model by

replacing the agent's incentive constraints with an appropriate (Bayesian) equilibrium. Second, the model could be extended to have an infinite time horizon. In this case our main theorem still holds assuming time-separable utility, discounting, and uniformly bounded felicity functions.

APPENDIX A

LEMMA 3. *If the mechanism $(\mathbf{x}^T, \mathbf{a}^T, \mathbf{p})$ is incentive compatible, the equilibrium payoff function of the agent, Π_0 , is Lipschitz continuous.*

PROOF. Throughout the proof, let K denote an integer such that the inequalities in [Assumption 0](#) are satisfied and, in addition, for all $t = 1, \dots, T$, $\tau < t$, and for all θ^t, a^t, x^t ,

$$\left| \frac{G_{t\theta_\tau}(\theta^t | \theta^{t-1}, a^{t-1}, x^{t-1})}{g_t(\theta^t | \theta^{t-1}, a^{t-1}, x^{t-1})} \right| < K.$$

First, we show that there exists a $\bar{K} \in \mathbb{N}$ such that $|\psi_{t\varepsilon_0}(\varepsilon^t, a^{t-1}, x^{t-1})| < \bar{K}$. For $t = 0$, $\psi_{0\varepsilon_0}(\varepsilon_0) = G_{0\varepsilon_0}^{-1}(\varepsilon_0) = 1/g_0(G_0^{-1}(\varepsilon_0)) < K$ by part (ii) of [Assumption 0](#). We proceed by induction and assume that $\psi_{\tau\varepsilon_0}(\varepsilon^\tau, a^{\tau-1}, x^{\tau-1}) < \bar{K}(\tau)$ for $\tau = 0, \dots, t-1$. Then

$$\begin{aligned} & |\psi_{t\varepsilon_0}(\varepsilon^t, a^{t-1}, x^{t-1})| \\ &= |G_{t\varepsilon_0}^{-1}(\varepsilon_t | \psi^{t-1}(\varepsilon^{t-1}, a^{t-2}, x^{t-2}), a^{t-1}, x^{t-1})| \\ &= \left| \frac{1}{g_t(\psi_t(\varepsilon^t, a^{t-1}, x^{t-1}) | \psi^{t-1}(\varepsilon^{t-1}, a^{t-2}, x^{t-2}), a^{t-1}, x^{t-1})} \right| \\ &\quad + \left| \sum_{\tau=0}^{\tau-1} G_{t\psi_\tau}^{-1}(\varepsilon_t | \psi^{t-1}(\varepsilon^{t-1}, a^{t-2}, x^{t-2}), a^{t-1}, x^{t-1}) \psi_{\tau\varepsilon_0}(\varepsilon^\tau, a^{\tau-1}, x^{\tau-1}) \right| \\ &\leq K + \max_{\tau \leq t} \bar{K}(\tau) \left| \sum_{\tau=0}^{\tau-1} G_{t\psi_\tau}^{-1}(\varepsilon_t | \psi^{t-1}(\varepsilon^{t-1}, a^{t-2}, x^{t-2}), a^{t-1}, x^{t-1}) \right|, \end{aligned}$$

where the inequality follows from the inductive hypothesis and part (ii) of [Assumption 0](#). However,

$$\begin{aligned} & K + \max_{\tau \leq t} \bar{K}(\tau) \left| \sum_{\tau=0}^{\tau-1} G_{t\psi_\tau}^{-1}(\varepsilon_t | \psi^{t-1}(\varepsilon^{t-1}, a^{t-2}, x^{t-2}), a^{t-1}, x^{t-1}) \right| \\ &= K + \max_{\tau \leq t} \bar{K}(\tau) \left| \sum_{\tau=0}^{\tau-1} G_{t\psi_\tau}(\psi_t(\varepsilon^t, a^{t-1}, x^{t-1}) | \psi^{t-1}(\varepsilon^{t-1}, a^{t-2}, x^{t-2}), a^{t-1}, x^{t-1}) \right| \\ &\quad \times \frac{1}{g_t(\psi_t(\varepsilon^t, a^{t-1}, x^{t-1}) | \psi^{t-1}(\varepsilon^{t-1}, a^{t-2}, x^{t-2}), a^{t-1}, x^{t-1})} \\ &\leq K + \max_{\tau \leq t} \bar{K}(\tau) K^2, \end{aligned}$$

by [Assumption 2](#). So we can conclude that $|\psi_{t\varepsilon_0}(\varepsilon^t, a^{t-1}, x^{t-1})| < K + \max_{\tau \leq t} \bar{K}(\tau) K^2$.

We are ready to prove that Π_0 is Lipschitz-continuous. Suppose that $\Pi_0(\varepsilon_0) \geq \Pi_0(\widehat{\varepsilon}_0)$. Let $\pi_0(\widehat{\varepsilon}_0, \varepsilon_0)$ denote the payoff of an agent whose initial type is $\widehat{\varepsilon}_0$, reports ε_0 , then reports truthfully afterward, and takes action $\widehat{\mathbf{a}}_t(\varepsilon^t, \widehat{\varepsilon}_0, s^{t-1})$ after history (ε^t, s^{t-1}) . Since the mechanism $(\mathbf{x}^T, \mathbf{a}^T, \mathbf{p})$ is incentive compatible, $\pi_0(\widehat{\varepsilon}_0, \varepsilon_0) < \Pi_0(\widehat{\varepsilon}_0)$ and hence,

$$\Pi_0(\varepsilon_0) - \Pi_0(\widehat{\varepsilon}_0) < \Pi_0(\varepsilon_0) - \pi_0(\widehat{\varepsilon}_0, \varepsilon_0).$$

So it is enough to prove that

$$|\Pi_0(\varepsilon_0) - \pi_0(\widehat{\varepsilon}_0, \varepsilon_0)| < K|\varepsilon_0 - \widehat{\varepsilon}_0|. \quad (18)$$

In addition,

$$\begin{aligned} \Pi_0(\varepsilon_0) - \pi_0(\widehat{\varepsilon}_0, \varepsilon_0) &= E[u(\varepsilon^T, \mathbf{a}^T(\varepsilon^T, s^{T-1}), s^T, \mathbf{x}^T(\varepsilon^T)) | \varepsilon_0] \\ &\quad - E[u(\varepsilon^T, \widehat{\mathbf{a}}^T(\varepsilon^T, \widehat{\varepsilon}_0, s^{T-1}), s^T, \mathbf{x}^T(\widehat{\varepsilon}_0, \varepsilon_{-0}^T, s^T)) | \varepsilon_0]. \end{aligned}$$

To establish (18) it is sufficient to show that the absolute value of the difference between the terms whose expectations are taken on the right-hand side of the previous equation is smaller than $K|\varepsilon_0 - \widehat{\varepsilon}_0|$. Note that

$$\begin{aligned} &u(\varepsilon^T, \mathbf{a}^T(\varepsilon^T, s^{T-1}), s^T, x^T) - u(\widehat{\varepsilon}_0, \varepsilon_{-0}^T, \widehat{\mathbf{a}}^T(\widehat{\varepsilon}_0, \varepsilon_{-0}^T, s^{T-1}), s^T, x^T) \\ &= \int_{\widehat{\varepsilon}_0}^{\varepsilon_0} u_{\varepsilon_0}(y, \varepsilon_{-0}^T, \mathbf{a}^T(y, \varepsilon_{-0}^T, s^{T-1}), s^T, x^T) \\ &\quad + \sum_{t=0}^T u_{a_t}(y, \varepsilon_{-0}^T, \mathbf{a}^T(y, \varepsilon_{-0}^T, s^{T-1}), s^T, x^T) \widehat{\mathbf{a}}_{t\widehat{\varepsilon}_0}(\varepsilon^t, y, s^{t-1}) dy. \end{aligned}$$

We will show that both terms on the right-hand side of the previous equation are bounded by a constant times $|\varepsilon_0 - \widehat{\varepsilon}_0|$. Note that

$$\begin{aligned} &\int_{\widehat{\varepsilon}_0}^{\varepsilon_0} u_{\varepsilon_0}(y, \varepsilon_{-0}^T, a^T, s^T, x^T) dy \\ &= \int_{\widehat{\varepsilon}_0}^{\varepsilon_0} \sum_{t=0}^T \widetilde{u}_{\theta_t}(\psi^T(y, \varepsilon_{-0}^T, a^{T-1}, x^{T-1}), a^T, s^T, x^T) \psi_{t\varepsilon_0}(y, \varepsilon_{-0}^T, a^{T-1}, x^{T-1}) dy \\ &\leq TK\overline{K}|\varepsilon_0 - \varepsilon| \end{aligned}$$

by part (i) of [Assumption 0](#) and since $\psi_{t\varepsilon_0}(\varepsilon^t) < \overline{K}$, as shown above. In addition,

$$\begin{aligned} &\int_{\widehat{\varepsilon}_0}^{\varepsilon_0} \sum_{t=0}^T u_{a_t}(y, \varepsilon_{-0}^T, \mathbf{a}^T(y, \varepsilon_{-0}^T, s^{T-1}), s^T, x^T) \widehat{\mathbf{a}}_{t\widehat{\varepsilon}_0}(\varepsilon^t, y, s^{t-1}) dy \\ &= \int_{\widehat{\varepsilon}_0}^{\varepsilon_0} \sum_{t=0}^T \widetilde{u}_{a_t}(\psi^T(y, \varepsilon_{-0}^T, \cdot), \mathbf{a}^T(y, \varepsilon_{-0}^T, s^{T-1}), s^T, x^T) \widehat{\mathbf{a}}_{t\widehat{\varepsilon}_0}(\varepsilon^t, y, s^{t-1}) dy. \end{aligned} \quad (19)$$

By the implicit function theorem,

$$\widehat{\mathbf{a}}_{t\widehat{\varepsilon}_0}(\varepsilon^t, y, s^{t-1}) = -\frac{f_{t\theta_t}(\psi^t(y, \varepsilon_{-0}^t, \cdot), \widehat{\mathbf{a}}_t(\varepsilon^t, y, s^{t-1}))}{f_{ta_t}(\psi^t(y, \varepsilon_{-0}^t, \cdot), \widehat{\mathbf{a}}_t(\varepsilon^t, y, s^{t-1}))} \psi_{t\widehat{\varepsilon}_0}(\varepsilon_{-0}^T, y, a^{T-1}, x^{T-1}),$$

which does not exceed $K\overline{K}$ by part (iii) of [Assumption 0](#) and the argument above showing that $|\psi_{t\widehat{\varepsilon}_0}| < \overline{K}$. Hence, (19) is smaller than

$$K\overline{K} \int_{\widehat{\varepsilon}_0}^{\varepsilon_0} \sum_{t=0}^T \widetilde{u}_{a_t}(\psi^T(y, \varepsilon_{-0}^T, \cdot), \mathbf{a}^T(y, \varepsilon_{-0}^T, s^{T-1}), s^T, x^T) dy \leq TK^2\overline{K}|\varepsilon_0 - \varepsilon|$$

by part (i) of [Assumption 0](#). □

PROOF OF LEMMA 1. Part (i) of [Assumption 2](#) is satisfied by definition. To see part (ii), notice that if $\overline{\mathbf{a}}_t(\theta_t, x^t)$ is interior, then

$$\begin{aligned} v_{t\theta_t}(\theta_t, x_t) &= \widetilde{u}_{\theta_t}(\theta_t, \overline{\mathbf{a}}_t(\theta_t, x^t), x^t) + \widetilde{u}_{ta_t}(\theta_t, \overline{\mathbf{a}}_t(\theta_t, x^t), x^t) \frac{\partial \overline{\mathbf{a}}_t(\theta_t, x^t)}{\partial \theta_t} \\ &= \widetilde{u}_{t\theta_t}(\theta_t, \overline{\mathbf{a}}_t(\theta_t, x^t), x^t) > 0, \end{aligned} \quad (20)$$

where the second equality follows from (8), and the inequality follows from part (ii) of [Assumption 2](#). If $\overline{\mathbf{a}}_t(\theta_t, x^t)$ is not interior then, generically,

$$v_{t\theta_t}(\theta_t, x_t) = \widetilde{u}_{t\theta_t}(\theta_t, \overline{\mathbf{a}}_t(\theta_t, x^t), x^t) > 0, \quad (21)$$

where the inequality again follows from part (ii) of [Assumption 2](#).

It remains to prove that v^t satisfies part (iii) of [Assumption 2](#). To simplify notation, we only prove this claim for the case when the contractible decision is unidimensional in each period, that is, $X_t \subset \mathbb{R}$ for all $t = 0, \dots, T$. Suppose first that $\overline{\mathbf{a}}_t(\theta_t, x^t)$ is interior. Note that for all $\tau \leq t$,

$$\begin{aligned} v_{t\theta_t x_\tau}(\theta_t, x_t) &= \widetilde{u}_{t\theta_t x_\tau}(\theta_t, \overline{\mathbf{a}}_t(\theta_t, x^t), x^t) + \widetilde{u}_{ta_t}(\theta_t, \overline{\mathbf{a}}_t(\theta_t, x^t), x^t) \frac{\partial \overline{\mathbf{a}}_t(\theta_t, x^t)}{\partial x_\tau} \\ &= \widetilde{u}_{t\theta_t x_\tau}(\theta_t, \overline{\mathbf{a}}_t(\theta_t, x^t), x^t) - \widetilde{u}_{t\theta_t a_t}(\theta_t, \overline{\mathbf{a}}_t(\theta_t, x^t), x^t) \frac{\widetilde{u}_{ta_t x_t}(\theta_t, \overline{\mathbf{a}}_t(\theta_t, x^t), x^t)}{\widetilde{u}_{ta_t^2}(\theta_t, \overline{\mathbf{a}}_t(\theta_t, x^t), x^t)}, \end{aligned}$$

where the first equality follows from (20) and the second one follows from (8) and the implicit function theorem. Note that $\widetilde{u}_{t\theta_t x_\tau}$, $\widetilde{u}_{t\theta_t a_t}$, and $\widetilde{u}_{ta_t x_t}$ are all nonnegative by part (iii) of [Assumption 2](#) and parts (ii) and (iii) of [Assumption 3](#). In addition, $\widetilde{u}_{ta_t^2}$ is negative by part (i) of [Assumption 3](#). Therefore, $v_{t\theta_t x_\tau}(\theta_t, x_t) \geq 0$. Suppose now that $\overline{\mathbf{a}}_t(\theta_t, x^t)$ is not interior. Then, for all $\tau \leq t$, generically,

$$v_{t\theta_t x_\tau}(\theta_t, x_t) = \widetilde{u}_{t\theta_t x_\tau}(\theta_t, \overline{\mathbf{a}}_t(\theta_t, x^t), x^t) \geq 0,$$

where the equality follows from (21) and the inequality follows from [Assumption 3\(ii\)](#). □

PROOF OF LEMMA 2. Part (i) of Assumption 2 is satisfied by definition. To see part (ii), notice that

$$w_{t\theta_t}(\theta_t, y_t, x^t) = \tilde{u}_{t\theta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t), x^t) + \tilde{u}_{ta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t), x^t) \frac{\partial \underline{\mathbf{a}}_t(\theta_t, y_t)}{\partial \theta_t}. \quad (22)$$

We apply the implicit function theorem for the identity $f_t(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t)) \equiv y_t$ to get

$$\frac{\partial \underline{\mathbf{a}}_t(\theta_t, y_t)}{\partial \theta_t} = -\frac{f_{t\theta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))}{f_{ta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))},$$

which is negative by part (ii) of Assumption 5. Since $\tilde{u}_{t\theta_t} > 0$ by part (ii) of Assumption 2 and $\tilde{u}_{ta_t} < 0$ by part (i) of Assumption 5, we conclude that w_t is strictly increasing in θ_t .

Next, we prove that w_t satisfies part (iii) of Assumption 2. First, we establish the single-crossing property with respect to θ_t and y_t . By (22),

$$\begin{aligned} w_{t\theta_t y_t}(\theta_t, y_t, x^t) &= \tilde{u}_{t\theta_t a_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t), x^t) \frac{\partial \underline{\mathbf{a}}_t(\theta_t, y_t)}{\partial y_t} \\ &\quad + \tilde{u}_{ta_t^2}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t), x^t) \frac{\partial \underline{\mathbf{a}}_t(\theta_t, y_t)}{\partial y_t} \frac{\partial \underline{\mathbf{a}}_t(\theta_t, y_t)}{\partial \theta_t} \\ &\quad + \tilde{u}_{ta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t), x^t) \frac{\partial^2 \underline{\mathbf{a}}_t(\theta_t, y_t)}{\partial \theta_t \partial y_t}. \end{aligned}$$

To sign $\partial \underline{\mathbf{a}}_t / \partial y_t$ and $\partial^2 \underline{\mathbf{a}}_t / \partial \theta_t \partial y_t$, we appeal to the implicit function theorem once again:

$$\begin{aligned} \frac{\partial \underline{\mathbf{a}}_t(\theta_t, y_t)}{\partial y_t} &= \frac{1}{f_{ta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))} \\ \frac{\partial^2 \underline{\mathbf{a}}_t(\theta_t, y_t)}{\partial \theta_t \partial y_t} &= \frac{f_{ta_t^2}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t)) \frac{f_{t\theta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))}{f_{ta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))} - f_{ta_t\theta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))}{f_{ta_t}^2(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))}. \end{aligned}$$

Therefore, $w_{t\theta_t y_t}(\theta_t, y_t, x^t)$ can be rewritten as

$$\begin{aligned} &\frac{\tilde{u}_{t\theta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t), x^t)}{f_{ta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))} + \tilde{u}_{ta_t^2}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t), x^t) \frac{-f_{t\theta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))}{f_{ta_t}^2(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))} \\ &\quad + \tilde{u}_{ta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t), x^t) \frac{f_{ta_t^2}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t)) \frac{f_{t\theta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))}{f_{ta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))} - f_{ta_t\theta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))}{f_{ta_t}^2(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))}. \end{aligned}$$

The first term is positive by part (ii) of Assumption 2 and part (ii) of Assumption 5. The second term is positive by part (i) of Assumption 3 and part (ii) of Assumption 5. The third term is positive by parts (i) and (iii) of Assumption 5. Therefore, we conclude that $w_{t\theta_t y_t} \geq 0$.

It remains to show that the single-crossing property in part (iii) of Assumption 2 also holds with respect to θ_t and x_τ for all $\tau \leq t$. To simplify notation, we only prove this

claim for the case when the contractible decision is unidimensional in each period, that is, $X_t \subset \mathbb{R}$ for all $t = 0, \dots, T$. By (22),

$$\begin{aligned} w_{t\theta_t x_t}(\theta_t, y_t, x^t) &= \tilde{u}_{t\theta_t x_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t), x^t) + \tilde{u}_{ta_t x_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t), x^t) \frac{\partial \underline{\mathbf{a}}_t(\theta_t, y_t)}{\partial \theta_t} \\ &= \tilde{u}_{t\theta_t x_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t), x^t) - \tilde{u}_{ta_t x_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t), x^t) \frac{f_{t\theta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))}{f_{ta_t}(\theta_t, \underline{\mathbf{a}}_t(\theta_t, y_t))}, \end{aligned}$$

which is positive by part (iv) of [Assumption 5](#). $\square$

PROOF OF PROPOSITION 4. Fix an increasing decision rule $(\tilde{\mathbf{x}}^T, \tilde{\mathbf{a}}^T)$ and a $\delta > 0$. Below, we construct a transfer rule, $\tilde{\mathbf{p}}$, such that $(\tilde{\mathbf{x}}^T, \tilde{\mathbf{a}}^T, \tilde{\mathbf{p}})$ satisfies (9). To this end, define the function $\tilde{\mathbf{y}}^T : \Theta^T \rightarrow Y^T$ such that $\tilde{\mathbf{y}}_t(\theta^t) = f_t(\theta_t, \tilde{\mathbf{a}}_t(\theta^t))$ for all t and θ^t . Since $\tilde{\mathbf{a}}_t$ is increasing in θ^t and f_t is strictly increasing in both θ_t and a_t (see part (ii) of [Assumption 5](#)), the function $\tilde{\mathbf{y}}_t$ is also increasing in θ^t . Therefore, by [Lemma 2](#) and [Proposition 2](#), the decision rule $(\tilde{\mathbf{x}}^T, \tilde{\mathbf{y}}^T)$ is implementable in a pure adverse selection model where the agent flow utilities are $\{w_t\}_{t=0}^T$. Let $\bar{\mathbf{p}} : \Theta^T \rightarrow \mathbb{R}$ denote a transfer rule that implements $(\tilde{\mathbf{x}}^T, \tilde{\mathbf{y}}^T)$.

Fix a $K \in \mathbb{N}$ such that $|\tilde{u}_{ta_t}| < K$ and $f_{ta_t} > 1/K$. By part (i) of [Assumption 0](#) and part (ii) of [Assumption 5](#), such a K exists. By Theorem 2 of [McAfee and Reny \(1992\)](#), for each $t = 0, \dots, T$, there exists a function $p_t : S_t \times Y_t \rightarrow \mathbb{R}$ such that $E_{s_t}(p_t(s_t, y_t) | f(\theta_t, a_t) = y_t) = 0$ and

$$E_{s_t}(p_t(s_t, y_t) | f(\theta_t, a_t) = y'_t) \geq K^2 |y_t - y'_t| - \frac{\delta}{T+1}. \quad (23)$$

Let us now define $\tilde{\mathbf{p}} : \Theta^T \times S^T \rightarrow \mathbb{R}$ by

$$\tilde{\mathbf{p}}(\theta^T, s^T) = \bar{\mathbf{p}}(\theta^T) + \sum_{t=0}^T p_t(s_t, \tilde{\mathbf{y}}_t(\theta^t)). \quad (24)$$

Next, we show that the agent cannot generate an excess payoff of δ by deviating from truth-telling and obedience in the mechanism $(\tilde{\mathbf{x}}^T, \tilde{\mathbf{a}}^T, \tilde{\mathbf{p}})$. First, note that the agent cannot benefit from making her strategy at time t contingent on the history of contractible signals, s^{t-1} , because her continuation payoff does not depend on these variables in the mechanism $(\tilde{\mathbf{x}}^T, \tilde{\mathbf{a}}^T, \tilde{\mathbf{p}})$. Therefore, we restrict attention to strategies that do not depend on past realizations of the contractible signal. Any such strategy induces a mapping from type profile to reports and actions in each period. Let $\rho_t(\theta^t)$ and $\alpha_t(\theta^t)$ denote the agent's report and action at time t , respectively, conditional on her type history θ^t . Let $\tilde{\alpha}_t(\theta^t)$ denote the solution of

$$f_t(\theta_t, a_t) = f_t(\rho_t(\theta^t), \tilde{\mathbf{a}}_t(\rho_t(\theta^t))) \quad [= \tilde{\mathbf{y}}_t(\rho_t(\theta^t))] \quad (25)$$

in a_t . In other words, $\tilde{\alpha}_t(\theta^t)$ is the agent's action that generates the same value of f_t conditional on θ^t as if the agent's true type was $\rho_t(\theta^t)$ and she took action $\tilde{\mathbf{a}}_t(\rho_t(\theta^t))$.

Then the expected payoff generated by (ρ^T, α^T) , conditional on θ_0 , is

$$\begin{aligned}
 & E_{\theta^T, s^T} \left[\sum_{t=0}^T \tilde{u}_t(\theta_t, \alpha_t(\theta^t), \tilde{\mathbf{x}}^t(\rho_t(\theta^t))) - \tilde{\mathbf{p}}(\rho^T(\theta^T), s^T) \middle| \theta_0 \right] \\
 &= E_{\theta^T} \left[\sum_{t=0}^T \tilde{u}_t(\theta_t, \tilde{\alpha}_t(\theta^t), \tilde{\mathbf{x}}^t(\rho_t(\theta^t))) - \bar{\mathbf{p}}(\rho^T(\theta^T)) \middle| \theta_0 \right] \\
 &\quad + \sum_{t=0}^T E_{\theta^T, s^T} [\tilde{u}_t(\theta_t, \alpha_t(\theta^t), \tilde{\mathbf{x}}^t(\rho_t(\theta^t))) - \tilde{u}_t(\theta_t, \tilde{\alpha}_t(\theta^t), \tilde{\mathbf{x}}^t(\rho_t(\theta^t))) \\
 &\quad - p_t(s_t, \tilde{\mathbf{y}}_t(\theta^t)) \middle| \theta_0],
 \end{aligned} \tag{26}$$

where the equality follows from (24).

We first consider the first term on the right-hand side of the previous equality. Note that

$$\begin{aligned}
 & E_{\theta^T} \left[\sum_{t=0}^T \tilde{u}_t(\theta_t, \tilde{\alpha}_t(\theta^t), \tilde{\mathbf{x}}^t(\rho_t(\theta^t))) - \bar{\mathbf{p}}(\rho^T(\theta^T)) \middle| \theta_0 \right] \\
 &= E_{\theta^T} \left[\sum_{t=0}^T w_t(\theta_t, \mathbf{y}_t(\rho_t(\theta^t)), \tilde{\mathbf{x}}^t(\rho_t(\theta^t))) - \bar{\mathbf{p}}(\rho^T(\theta^T)) \middle| \theta_0 \right] \\
 &\leq E_{\theta^T} \left[\sum_{t=0}^T w_t(\theta_t, \mathbf{y}_t(\theta^t), \tilde{\mathbf{x}}^t(\theta^t)) - \bar{\mathbf{p}}(\theta^T) \middle| \theta_0 \right] \\
 &= E_{\theta^T} \left[\sum_{t=0}^T \tilde{u}_t(\theta_t, \tilde{\mathbf{a}}_t(\theta^t), \tilde{\mathbf{x}}^t(\theta^t)) - \bar{\mathbf{p}}(\theta^T) \middle| \theta_0 \right],
 \end{aligned} \tag{27}$$

where the inequality follows from the assumption that the transfer rule $\bar{\mathbf{p}}$ implements $(\tilde{\mathbf{x}}^T, \tilde{\mathbf{y}}^T)$ if the flow utilities are $\{w^t\}_{t=0}^T$. Also note that

$$\begin{aligned}
 & \tilde{u}_t(\theta_t, \alpha_t(\theta^t), \tilde{\mathbf{x}}^t(\rho_t(\theta^t))) - \tilde{u}_t(\theta_t, \tilde{\alpha}_t(\theta^t), \tilde{\mathbf{x}}^t(\rho_t(\theta^t))) - E_{s^T} [p_t(s_t, \tilde{\mathbf{y}}_t(\rho_t(\theta^T))) \middle| \theta_t, \alpha_t(\theta^t)] \\
 &\leq K |\tilde{\alpha}_t(\theta^t) - \alpha_t(\theta^t)| - E_{s^T} [p_t(s_t, \tilde{\mathbf{y}}_t(\rho_t(\theta^T))) \middle| \theta_t, \alpha_t(\theta^t)] \\
 &\leq K^2 |f_t(\theta_t, \tilde{\alpha}_t(\theta^t)) - f_t(\theta_t, \alpha_t(\theta^t))| - E_{s^T} [p_t(s_t, \tilde{\mathbf{y}}_t(\rho_t(\theta^T))) \middle| \theta_t, \alpha_t(\theta^t)] \\
 &= K^2 |\tilde{\mathbf{y}}_t(\rho_t(\theta^T)) - f_t(\rho_t(\theta^t), \alpha_t(\theta^t))| - E_{s^T} [p_t(s_t, \tilde{\mathbf{y}}_t(\rho_t(\theta^T))) \middle| \theta_t, \alpha_t(\theta^t)] \\
 &\leq \frac{\delta}{T+1},
 \end{aligned}$$

where the first and second inequalities follow from $|\tilde{u}_{ta_t}| < K$ and $f_{ta_t} > 1/K$, the equality follows from (25), and the last inequality follows from (23). Summing up these inequalities for $t = 0, \dots, T$ and taking expectation with respect to θ^T yields

$$\sum_{t=0}^T E_{\theta^T, s^T} [\tilde{u}_t(\theta_t, \alpha_t(\theta^t), \tilde{\mathbf{x}}^t(\rho_t(\theta^t))) - \tilde{u}_t(\theta_t, \tilde{\mathbf{a}}_t(\rho_t(\theta^t)), \tilde{\mathbf{x}}^t(\rho_t(\theta^t))) - p_t(s_t, \tilde{\mathbf{y}}_t(\theta^t)) | \theta_0] \leq \delta. \quad (28)$$

Therefore, plugging (27) and (28) into (26) we get that

$$\begin{aligned} E_{\theta^T, s^T} \left[\sum_{t=0}^T \tilde{u}_t(\theta_t, \alpha_t(\theta^t), \tilde{\mathbf{x}}^t(\rho_t(\theta^t))) - \tilde{\mathbf{p}}(\rho^T(\theta^T), s^T) \middle| \theta_0 \right] \\ \leq E_{\theta^T} \left[\sum_{t=0}^T \tilde{u}_t(\theta_t, \tilde{\mathbf{a}}_t(\theta^t), \tilde{\mathbf{x}}^t(\theta^t)) - \tilde{\mathbf{p}}(\theta^T) \middle| \theta_0 \right] + \delta \\ = E_{\theta^T, s^T} \left[\sum_{t=0}^T \tilde{u}_t(\theta_t, \tilde{\mathbf{a}}_t(\theta^t), \tilde{\mathbf{x}}^t(\theta^t)) - \tilde{\mathbf{p}}(\rho^T(\theta^T), s^T) \middle| \theta_0 \right] + \delta, \end{aligned}$$

where the equality follows from $E_{s_t} [p_t(s_t, y_t) | f(\theta_t, a_t) = y_t] = 0$. This implies that the agent cannot gain more than δ by deviating from truth-telling and obedience in the mechanism $(\tilde{\mathbf{x}}^T, \tilde{\mathbf{a}}^T, \tilde{\mathbf{p}})$. $\square$

REFERENCES

- Baron, David P. and David Besanko (1984), “Regulation and information in a continuing relationship.” *Information Economics and Policy*, 1, 267–302. [111, 112, 117]
- Battaglini, Marco and Rohit Lamba (2014), “Optimal dynamic contracting: The first-order approach and beyond.” Unpublished, SSRN 2172414. [110]
- Bergemann, Dirk and Juuso Välimäki (2002), “Information acquisition and efficient mechanism design.” *Econometrica*, 70, 1007–1033. [129]
- Courty, Pascal and Hao Li (2000), “Sequential screening.” *Review of Economic Studies*, 67, 697–717. [111, 117]
- Eső, Péter and Balázs Szentes (2007a), “Optimal information disclosure in auctions and the handicap auction.” *Review of Economic Studies*, 74, 705–731. [111, 117, 122]
- Eső, Péter and Balázs Szentes (2007b), “The price of advice.” *RAND Journal of Economics*, 38, 863–880. [111]
- Farhi, Emmanuel and Iván Werning (2013), “Insurance and taxation over the life cycle.” *Review of Economic Studies*, 80, 596–635. [111]
- Garrett, Daniel F. and Alessandro Pavan (2012), “Managerial turnover in a changing world.” *Journal of Political Economy*, 120, 879–925. [111, 112, 116, 117, 130]

Golosov, Mikhail, Maxim Troshkin, and Aleh Tsyvinski (2011), “Optimal dynamic taxes.” NBER Working Paper 17642. [111]

Kapička, Marek (2013), “Efficient allocations in dynamic private information economies with persistent shocks: A first-order approach.” *Review of Economic Studies*, 80, 1027–1054. [111]

Krähmer, Daniel and Roland Strausz (2015a), “Dynamics of mechanism design.” In *An Introduction to the Theory of Mechanism Design* (Tilman Börgers, ed.), 204–234, Oxford University Press, New York. [111]

Krähmer, Daniel and Roland Strausz (2015b), “Ex post information rents in sequential screening.” *Games and Economic Behavior*, 90, 257–273. [117]

Laffont, Jean-Jacques and Jean Tirole (1986), “Using cost observation to regulate firms.” *Journal of Political Economy*, 94, 614–641. [112]

McAfee, R. Preston and Philip J. Reny (1992), “Correlated information and mechanism design.” *Econometrica*, 60, 395–421. [121, 124, 126, 136]

Milgrom, Paul R. and Ilya Segal (2002), “Envelope theorems for arbitrary choice sets.” *Econometrica*, 70, 583–601. [118]

Pavan, Alessandro, Ilya Segal, and Juuso Toikka (2014), “Dynamic mechanism design: A Myersonian approach.” *Econometrica*, 82, 601–653. [111, 112, 116, 117, 122]

Co-editor Johannes Hörner handled this manuscript.

Manuscript received 17 March, 2015; final version accepted 18 October, 2015; available online 23 December, 2015.