

SCREENING FOR BREAKTHROUGHS*

Gregorio Curello

Ludvig Sinander

Abstract

How best to incentivise *prompt* disclosure? We study this question in a general model in which a technological breakthrough occurs at an uncertain time and is privately observed by an agent, and a principal must incentivise disclosure via her control of a payoff-relevant physical allocation. We uncover a deadline structure of optimal mechanisms: they have a simple deadline form in an important special case, and a graduated deadline structure in general. We apply our results to the design of unemployment insurance schemes.

1 Introduction

Society advances by finding better ways of doing things. When such a technological breakthrough occurs, it frequently becomes known only to certain individuals with particular expertise. Only if such individuals share their knowledge promptly can the promise of progress be unlocked.

The resulting need to incentivise prompt disclosure engenders a screening problem in which the agent's private information is about *when*, rather than about *what*. We call this *screening for breakthroughs*.

*We are deeply grateful to Eddie Dekel for many detailed comments. Sinander thanks him, Alessandro Pavan and Bruno Strulovici for their guidance and support. We have profited from comments from anonymous referees, Dilip Abreu, Miguel Ballester, Lori Beaman, Iván Canay, Gabe Carroll, Sylvain Chassang, Janet Currie, Théo Durandard, Piotr Dworczak, Andrew Ellis, Jeff Ely, Alex Frug, George Georgiadis, Brett Green, Faruk Gül, Yingni Guo, Marina Halac, Oliver Hart, Ian Jewitt, Shengwu Li, Alessandro Lizzeri, Guido Lorenzoni, Erik Madsen, Thomas Mariotti, David Martimort, Eric Maskin, John Moore, Matt Notowidigdo, Wojciech Olszewski, Harry Pei, Wolfgang Pesendorfer, Dan Quigley, Debraj Ray, Ronny Razin, Wolfgang Ridinger, Marciano Siniscalchi, Juuso Toikka, Arne Uhlenhorff, Can Ürgün, Chris Wallace, Asher Wolinsky, Leeat Yariv and audiences at Bonn, Harvard, LSE, Lund, Northwestern, NYU, Oxford, Princeton, *Seminars in Economic Theory*, Toulouse, and the 2021 *REStud* Tour. Curello acknowledges support from the German Research Foundation (DFG) through CRC TR 224 (Project B02).

The need to screen for breakthroughs is widespread. One example is the much-discussed problem of talent-hoarding in organisations (see Hägele, 2024). The manager of a team is well-placed to know when one of her subordinates acquires a skill. When this happens, headquarters may wish to re-assign the worker to a new role better-suited to her abilities. Managers, however, have a documented tendency to want to hold on to their workers. Careful design is thus needed to incentivise prompt disclosure.

Another example is unemployment insurance: since unemployed workers are typically privately informed about when they receive a job offer, benefits must be designed with a view to incentivising them to start work promptly. A third example concerns technical innovations that reduce firms' greenhouse-gas emissions, at the price of raising production costs.¹ Only with suitable regulation will firms which discover such innovations choose to adopt them.

In this paper, we study the general problem of screening for breakthroughs. We introduce a model in which an agent privately observes when a new productive technology arrives. This breakthrough expands utility possibilities for the agent and principal, but generates a conflict of interest between them. The agent decides whether and when to disclose the breakthrough, and the principal controls a payoff-relevant physical allocation over time. Our model deliberately focusses on the screening-for-breakthroughs problem, excluding well-understood frictions such as the need to incentivise the agent to exert unobservable effort. It can be shown that adding such a moral-hazard friction to the model does not affect our results.²

We ask how the principal can best incentivise prompt disclosure of the breakthrough. Our answer uncovers a deadline structure of optimal mechanisms: only simple *deadline mechanisms* are optimal in an important special case, while a graduated deadline structure characterises optimal incentives in general. We apply these insights to the design of unemployment insurance schemes.

1.1 Overview of model and results

A breakthrough occurs at a random time, making available a new technology that expands utility possibilities for an agent and a principal. There is a conflict of interest: were the principal to operate the old and new technologies in her own interest, the agent would be better off under the old one. The agent privately observes when the breakthrough occurs, and (verifiably) discloses

¹Such innovations are expected to account for the bulk of abatement in the cement industry, currently the source of about 7% of all CO₂ emissions (Czigler et al., 2020).

²We omit this extension, but it may be found in the working paper (Curello & Sinander, 2025).

it at a time of her choosing. The principal controls a physical allocation that determines the agent's utility over time. (The description of a physical allocation may include a specification of monetary payments to the agent; in this case, the conflict-of-interest assumption requires that the agent be protected by (at least a degree of) limited liability prior to disclosure.)

To focus on the robust qualitative features of optimal screening, we study *undominated* mechanisms, meaning those such that no alternative mechanism is weakly better for the principal under any arrival distribution of the breakthrough and strictly better under some distribution. We further describe, for any given breakthrough distribution, the principal's optimal choice among undominated mechanisms.

Toward our deadline characterisation, we first study how undominated mechanisms incentivise the agent. We show that the agent should be indifferent at all times between prompt and delayed disclosure (Proposition 0). This is despite the fact that the standard argument fails: were the agent strictly to prefer prompt to delayed disclosure, then lowering the agent's post-disclosure utility would *not* necessarily benefit the principal.

We then elucidate the deadline structure of undominated mechanisms when the pre-breakthrough technology's utility possibilities have an affine shape. Theorem 1 asserts that in this case, all undominated mechanisms belong to a small class of simple *deadline mechanisms*. Absent disclosure, these mechanisms give the agent a Pareto-efficient utility u^0 before a deadline, and an inefficiently low utility u^* afterwards.³ The proof of Theorem 1 argues (loosely) that any mechanism may be improved by *front-loading* the agent's pre-disclosure utility, making it higher early and lower late while preserving its total discounted value. We further characterise the principal's optimal choice of deadline as a function of the breakthrough distribution (Proposition 2).

Outside of the affine case, optimal mechanisms exhibit a graduated deadline structure (Theorem 2): absent disclosure, the agent's utility still starts at the efficient level u^0 and declines monotonically toward the inefficiently low level u^* , but the transition may be gradual. For any given breakthrough distribution, we describe the optimal transition (Proposition 3).

We then apply our results to the design of unemployment insurance schemes. An unemployed worker (agent) receives a job offer at a random time, and chooses whether to accept, and if so how soon to start. Offers are private, but the state (principal) observes when the worker starts a job. The state controls unemployment benefits and income taxes, and cares both about the worker's welfare and net tax revenue.

³ u^0 and u^* are functions of the technologies, so the deadline is the only free parameter.

Many countries, such as Germany and France, pay a generous unemployment benefit until a deadline, and provide only a low benefit to those remaining unemployed beyond this deadline. Our results provide a potential rationale for such deadline schemes: they are approximately optimal provided that either (a) the worker’s consumption utility has limited curvature, or (b) tax revenue is comparatively unimportant for social welfare. Conversely, our analysis suggests that where neither (a) nor (b) is satisfied, substantial welfare gains could be achieved by tapering benefits gradually, as in Italy.

1.2 Related literature

This paper belongs to the literature on incentive design for a proposing agent, initiated by Armstrong and Vickers (2010).⁴ In their (static) model, the agent privately observes which physical allocations are available, then proposes one (or several). The key assumptions are that

- (a) the agent can propose only available allocations, and that
- (b) the principal can implement only proposed allocations.

Our dynamic problem shares these key features: the new technology (a) can only be disclosed (proposed) once available, and (b) can be utilised by the principal only once disclosed.

Bird and Frug (2019) study a different dynamic environment with features (a) and (b). Payoffs are simple: there is an allocation α preferred by the principal and a default allocation favoured by the agent,⁵ and the principal can furthermore reward the agent at a linear cost. In each period, the agent privately observes whether α is available; it can (a) be disclosed only if available, and (b) be implemented only if disclosed. Were rewards unrestricted, α could be implemented whenever available by rewarding the agent just enough to induce disclosure. (And this is optimal; thus there is no conflict of interest in our sense.) The authors instead subject promised rewards to a dynamic budget constraint,⁶ and study how the budget should be spent over time. By comparison, we allow for general payoffs (technologies) and impose no dynamic constraints, focussing instead on a conflict of interest.

⁴See also Nocke and Whinston (2013) and Guo and Shmaya (2023). Our account of the literature follows the latter authors’ insightful discussion. The literature has precedents in applied work on corporate finance (Berkovitch & Israel, 2004) and antitrust (Lyons, 2003).

⁵There is an extension to multiple allocations α ; little changes.

⁶They assume in particular that the agent can be rewarded only using exogenous reward ‘opportunities’, which arrive randomly over time; but nothing changes if rewards take other forms, e.g. (flow) monetary payments subject to a per-period cap.

Feature (a) means that the agent’s disclosures are verifiable, a possibility first studied by Viscusi (1978), Grossman and Hart (1980), Milgrom (1981), and Grossman (1981). A strand of the subsequent literature examines the role of commitment in static models,⁷ while another studies the timing of disclosure absent commitment;⁸ our environment features both commitment and dynamics.⁹ These models lack property (b): the agent cannot constrain the principal.

More distantly related is the large literature on dynamic adverse-selection models with cheap-talk communication (contrast with (a)) and no scope for the agent to constrain the principal’s choice of allocation (contrast with (b)). The strand on dynamic ‘delegation’ allows for non-transferable utility, as we do;¹⁰ otherwise the literature tends to focus on monetary transfers.¹¹ A recent strand examines models which, like ours, feature private information about *when*, rather than about *what*. For example, Green and Taylor (2016) show how moral hazard may be mitigated by conditioning pay and termination on cheap-talk ‘progress reports’.¹² In their model, the agent privately observes the arrival of a signal which indicates that project completion is within reach (given enough effort). Completion is observable. There is no conflict of interest in our sense; instead, the challenge is to incentivise unobservable (completion-hastening) effort. (Absent this moral hazard, the principal would have no reason to elicit the signal.) Relatedly, Madsen (2022) studies how cheap-talk progress reports may be elicited by conditioning pay and termination on a contractible signal. In his model, the agent privately observes when a project ‘expires’, and the principal decides when to terminate the project. The principal (agent) prefers termination close to expiry (as late as possible).

⁷Particularly Glazer and Rubinstein (2004, 2006), Sher (2011), Hart, Kremer, and Perry (2017), and Ben-Porath, Dekel, and Lipman (2019).

⁸See Dye and Sridhar (1995), Acharya, DeMarzo, and Kremer (2011), Guttman, Kremer, and Skrzypacz (2014), Campbell, Ederer, and Spinnewijn (2014), and Curello (2023a, 2023b). The last three papers feature ‘breakthroughs’, but these engender no conflict of interest in our sense; the incentive problem is instead that of deterring shirking.

⁹So does recent work on revenue management, where a firm contracts with customers who arrive unobservably over time and choose when verifiably to reveal themselves; see Pai and Vohra (2013), Board and Skrzypacz (2016), Mierendorff (2016), Garrett (2016, 2017), Gershkov, Moldovanu, and Strack (2018), and Dilmé and Li (2019).

¹⁰See Jackson and Sonnenschein (2007), Matsushima, Miyazaki, and Yagi (2010), Frankel (2016), Guo (2016), Li, Matouschek, and Powell (2017), Lipnowski and Ramos (2020), Guo and Hörner (2020), and de Clippel, Eliaz, Fershtman, and Rozen (2021).

¹¹E.g. Roberts (1982), Baron and Besanko (1984), Courty and Li (2000), Battaglini (2005), Eső and Szentes (2007a, 2007b), Board (2007), and Pavan, Segal, and Toikka (2014).

¹²See also Feng, Taylor, Westerfield, and Zhang (2024).

Crucially, there is a noisy contractible signal of expiry.¹³ Both of these papers use the term ‘deadline’, as we do, but mean quite different things by it.¹⁴

1.3 Roadmap

We introduce the model in the next section, then formulate the principal’s problem in §3. In §4, we show that undominated mechanisms incentivise the agent by keeping her always indifferent. We then describe the deadline structure of optimal mechanisms (§5 and §6). In §7, we apply our results to the design of unemployment insurance schemes.

2 Model

There is an agent and a principal, whose utilities are denoted by $u \in [0, \infty)$ and $v \in [-\infty, \infty)$, respectively. A frontier $F^0 : [0, \infty) \rightarrow [-\infty, \infty)$ describes utility possibilities: $F^0(u)$ is the highest utility that the principal can attain subject to giving the agent utility u . We assume that F^0 is concave and upper semi-continuous, that it has a unique peak $u^0 > 0$ (namely, $F^0(u^0) > F^0(u)$ for every $u \neq u^0$), and that it is finite on $(0, u^0]$. Such a frontier is depicted in Figure 1.

Time $t \in \mathbf{R}_+$ is continuous. The principal controls the agent’s flow utility u (and thus her own utility $F^0(u)$) over time, and is able to commit.

We interpret this abstract description of utility possibilities in the standard fashion: there is an (unmodelled) set of feasible physical allocations over which the agent and principal have preferences, and the principal decides which allocation prevails in each period. She thus effectively controls the agent’s flow utility. We illustrate and interpret further in §2.1 below.

At a random time τ , a *breakthrough* occurs: a new technology becomes available which expands the utility possibility frontier to $F^1 \geq F^0$. The new frontier is likewise concave and upper semi-continuous, with a unique peak denoted by $u^1 \geq 0$. (Note that we allow for the possibility that $u^1 = 0$, in which case F^1 is decreasing.) The breakthrough engenders a conflict of interest: the new frontier peaks at a strictly lower agent utility ($u^1 < u^0$),

¹³If there were no contractible signal, then non-trivial screening would be impossible, since the agent’s preferences are the same whatever her type (expiry date).

¹⁴Deterministic hard deadlines, as in our result, appear only in the benchmark case of Green and Taylor in which there is no signal (a case unrelated to our model and Madsen’s). In Green and Taylor, ‘(soft) deadline’ means a time after which termination may randomly occur if the agent has not yet reported the signal’s arrival. Madsen uses ‘(soft) deadline’ to mean that termination depends on the realisation of the contractible signal.

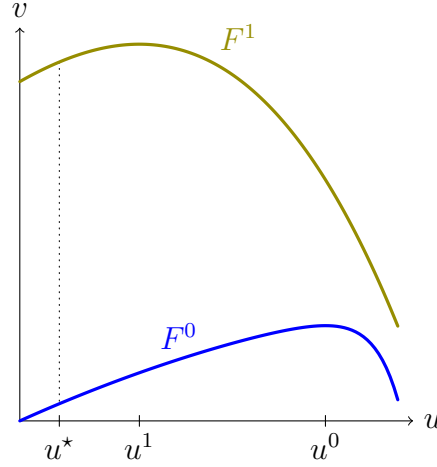


Figure 1: Utility possibility frontiers. The new technology expands utility possibilities ($F^1 \geq F^0$), but creates a conflict of interest ($u^1 < u^0$). u^* denotes the rightmost point to the left of u^0 at which F^0, F^1 have equal slopes.

so that the breakthrough would hurt the agent were the principal to operate both technologies in her own interest. This is illustrated in Figure 1.

The breakthrough is observed only by the agent. At any time $t \geq \tau$ after the breakthrough, she can verifiably disclose to the principal that it has occurred. (That is, she can *prove* that the new technology is available.) The new technology can be used only once its availability has been disclosed.

The agent and principal discount their flow payoffs at rate $r > 0$ and have expected-utility preferences, so that their respective payoffs from random flow utilities $t \mapsto x_t$ and $t \mapsto y_t$ are

$$\mathbf{E} \left(r \int_0^\infty e^{-rt} x_t dt \right) \quad \text{and} \quad \mathbf{E} \left(r \int_0^\infty e^{-rt} y_t dt \right).$$

The random time τ at which the breakthrough occurs is distributed according to an arbitrary cumulative distribution function G .

We write u^* for the rightmost $u \in [0, u^0]$ at which the old and new frontiers F^0, F^1 have equal slope (in the sense of sharing a supergradient—see Rockafellar (1970, part V)), and let $u^* := 0$ in case no such $u \in [0, u^0]$ exist. This utility level will feature prominently in our analysis. To avoid trivialities, we impose the weak genericity assumption that u^* is a strict local maximum of $F^1 - F^0$. Note that $u^* \leq u^1 < u^0$.

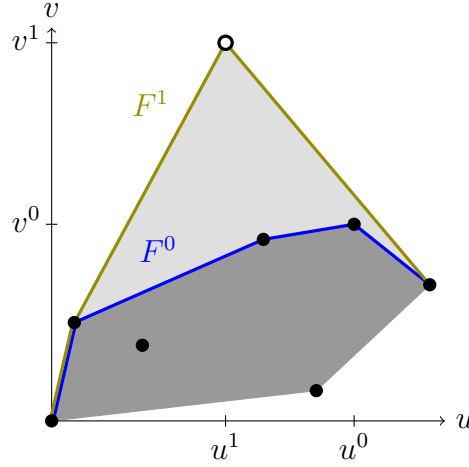


Figure 2: Finitely many allocations: the old ($\bullet$), the new ($\circ$), and utility possibilities (grey). Here $u^* = u^1$.

2.1 Interpreting the frontiers

In the simplest applications, there are finitely many (old) allocations, and the agent privately observes when a single new allocation becomes available. For example, a manager may observe when a worker on her team acquires a skill, or a firm may discover an emissions-reducing innovation. Each allocation provides some utilities (u, v) to the agent and principal, which may be plotted as in Figure 2. The utility possibility set is the convex hull of these profiles,¹⁵ and the frontier F^0 is its upper boundary. The agent privately observes when a new allocation (u^1, v^1) becomes available. The principal likes the new allocation better than any other, whereas the agent prefers the principal's favourite old allocation (u^0, v^0) . Thus utility possibilities expand, but there is a conflict of interest.

Example 1. The simplest formalisation of the talent-hoarding story from the introduction is as follows. A worker belongs to a team in an organisation. Her productivity on the team is $v^0 > 0$, while her productivity outside of the team is strictly lower, normalised to zero. At some uncertain time, she acquires a skill that can be exercised only outside of her current team, at productivity $v^1 > v^0$. (This could be the skill to manage a team of her own, for example.) Headquarters (the principal) cares about output, while the worker's manager (the agent) has a pure empire-building motive: her payoff is $u = 1$ if the worker is on her team and $u = 0$ otherwise. In this case, the

¹⁵In-between profiles are achieved by rapidly switching back and forth (or randomising).

frontiers are given by $F^0(u) = uv^0$ and $F^1(u) = (1 - u)v^1 + uv^0$ for each $u \in [0, 1]$.¹⁶ These frontiers satisfy our model assumptions; in particular, the conflict-of-interest assumption holds since $u^1 = 0 < 1 = u^0$.

Richer applications feature (infinitely) many allocations. In our application to unemployment insurance (§7), for example, an allocation specifies the worker’s consumption and (if she is employed) her labour supply.

Our abstract treatment of allocations allows for a broad range of applications. Allocations may be multi-dimensional, for example, with some dimensions corresponding to observable actions taken by the agent. (The principal controls these by issuing action recommendations, backed by the threat of giving the agent zero utility forever unless she complies.) One dimension of the allocation may describe monetary payments to the agent; we discuss this possibility in §2.2 below.

Rich downstream interactions between the principal and agent can be accommodated by re-interpreting the frontier F^1 in lifetime terms, so that $F^1(u)$ is the principal’s continuation utility from the post-disclosure interaction when she is constrained to provide the agent with a continuation utility of u .¹⁷ The post-disclosure interaction could be one of contracting under (rich, possibly dynamic) moral hazard, for example: that yields a frontier F^1 which satisfies our shape assumptions (see e.g. Sannikov, 2008, Figure 1).

2.2 Discussion of the assumptions

Two of our assumptions are economically substantive. First, the agent privately observes a technological breakthrough, but cannot utilise the new technology without the principal’s knowledge. Many economic environments have this feature: in unemployment insurance, for instance, the state observes the worker’s employment status (from e.g. tax records).

Secondly, there is a conflict of interest, captured by $u^1 < u^0$. Such conflicts arise naturally in applications: in unemployment insurance, for example, the state (principal) would like an employed worker (agent) to work and pay taxes, but the worker would rather not. Absent a conflict of interest, the principal can attain first-best (see Remark 1 below).

Many of the remaining model assumptions are innocuous, as we next briefly relate. For more details, see the working paper (Curello & Sinander, 2025).

¹⁶And $F^0(u) = F^1(u) = -\infty$ for all $u \in (1, \infty)$, since $u > 1$ is impossible.

¹⁷The legitimacy of this re-interpretation is formally established in §3.1 below. Note that the pre-breakthrough frontier F^0 cannot be re-interpreted in this ‘lifetime’ fashion.

Utility possibilities. The assumption that $F^1 \geq F^0$ is without loss of generality (the old technology remains available after the breakthrough, so the principal can still attain utility $\geq F^0(u)$ while giving the agent utility u , for every $u \in [0, \infty)$). The assumption that the frontiers are concave is likewise without loss: if one of them were not, then the principal could get arbitrarily close to any point on its concave upper envelope by rapidly switching back and forth between agent utility levels. Upper semi-continuity is similarly innocuous. The stipulation that u^* is a strict local maximum of $F^1 - F^0$ essentially just rules out a saddle point, and is anyway dispensable.

Not every agent utility $u \in [0, \infty)$ need be feasible: if no physical allocation provides utility u , then we let $F^j(u) := -\infty$, ensuring that u is never chosen by the principal. Our assumption that F^0 is finite on $(0, u^0]$ is without loss.

We have required the agent's flow utility u to be non-negative, meaning that there is a bound (normalised to zero) on how much misery the principal can inflict on the agent. This assumption may be replaced with a participation constraint without affecting our results.

Distribution. The distribution G of the breakthrough time is unrestricted: it can have atoms, for example, and need not have full support. It can be shown that our results extend to the case in which G is endogenously generated by the agent's unobservable exertion of costly effort.

Uncertain technology. Our analysis applies unchanged if the new frontier F^1 is random, provided the agent does not have private information about its realisation.

Cheap talk. Nothing changes if the agent's disclosures are non-verifiable, provided the principal observes her own payoffs in real time, since she can then verify cheap-talk reports at negligible cost.¹⁸

(Non-)transferable utility. The frontiers F^0, F^1 can encode monetary transfers between the principal and agent; our model assumptions restrict such transfers only by requiring that *before* the breakthrough, the agent is protected by (at least a degree of) limited liability. In detail, write $\tilde{F}^0, \tilde{F}^1$ for the frontiers describing utility possibilities *absent* monetary transfers. If the principal gives the agent gross utility $\tilde{u} \in [0, \infty)$ and pays her $w \in \mathbf{R}$, then net flow utilities are $\tilde{u} + w$ for the agent and $\tilde{F}^j(\tilde{u}) - w$ for the principal

¹⁸Following a report, the principal can provide utility u^1 for a short time, earning $F^1(u^1)$ if the breakthrough really did occur and $F^0(u^1) < F^1(u^1)$ if not.

when technology $j \in \{0, 1\}$ is used. Any constraints on payments, such as limited liability, are captured by constraint sets $W^0 \subseteq \mathbf{R}$ before disclosure and $W^1 \subseteq \mathbf{R}$ after disclosure. For $j \in \{0, 1\}$, the utility possibility frontier F^j equals the concave upper semi-continuous upper envelope of

$$u \mapsto \sup_{w \in W^j} \left\{ \tilde{F}^j(\tilde{u}) - w : \tilde{u} + w = u \right\}.$$

Assume that $\tilde{F}^0, \tilde{F}^1$ satisfy the model assumptions. Then F^0, F^1 also satisfy all model assumptions, *except* possibly for the conflict-of-interest assumption $u^1 < u^0$. What is needed for $u^1 < u^0$ to hold is that the agent be protected by a degree of limited liability *before* the breakthrough, i.e. $\inf W_0 \geq -k$ for some $k \in \mathbf{R}_+$; in particular, this condition with $k = 0$ is sufficient, and it is necessary for this condition to hold with *some* $k \geq 0$.¹⁹ The model assumptions imply no restrictions on *post*-disclosure payments W^1 .

2.3 Mechanisms and incentive-compatibility

A *mechanism* specifies, for each period $t \in \mathbf{R}_+$, the flow utility x_t^0 that the agent enjoys at t if she has not yet disclosed, as well as the continuation utility X_t^1 that she earns by disclosing at t . Formally, a mechanism is a pair (x^0, X^1) , where $x^0 : \mathbf{R}_+ \rightarrow \mathbf{R}_+$ and $X^1 : \mathbf{R}_+ \rightarrow [0, \infty]$ are Lebesgue-measurable. We call x^0 the *pre-disclosure flow*, and X^1 the *disclosure reward*.

(Our notation uses lowercase for flows and uppercase for stocks: flow utilities are $x_t \in [0, \infty)$, while continuation payoffs are $X_t \in [0, \infty]$. As usual, ‘ x ’ and ‘ X ’ denote the functions $t \mapsto x_t$ and $t \mapsto X_t$, respectively.)

Note that the description of a mechanism does not specify what utility flow $s \mapsto x_s^{1,t}$ the agent enjoys after disclosing at t , only its present value

$$X_t^1 = r \int_t^\infty e^{-r(s-t)} x_s^{1,t} ds$$

(which may be equal to ∞). Nor does the definition specify which technology is used when both are available. These omissions do not matter for the agent’s incentives, so we shall address them when we formulate the principal’s problem (next section).

A mechanism is *incentive-compatible* (‘*IC*’) iff the agent prefers disclosing promptly to (a) disclosing with a delay or (b) never disclosing. Formally:

Definition 1. A mechanism (x^0, X^1) is *incentive-compatible* (‘*IC*’) iff for every period $t \in \mathbf{R}_+$,

¹⁹Write $\tilde{u}^0, \tilde{u}^1$ for the peaks of $\tilde{F}^0, \tilde{F}^1$, and note that $u^1 \leq \tilde{u}^1 < \tilde{u}^0 \geq u^0$. If $\inf W_0 \geq 0$ then $u^0 = \tilde{u}^0$, so $u^1 \leq \tilde{u}^1 < \tilde{u}^0 = u^0$. If $\inf W_0 < -k$ for every $k \in \mathbf{R}_+$, then $u^0 = 0 \leq u^1$.

- (a) $X_t^1 \geq r \int_t^{t+d} e^{-r(s-t)} x_s^0 ds + e^{-rd} X_{t+d}^1$ for every $d > 0$, and
- (b) $X_t^1 \geq r \int_t^\infty e^{-r(s-t)} x_s^0 ds$.

By a revelation principle à la Bull and Watson (2007), we may restrict attention to incentive-compatible mechanisms (for details, see the working paper (Curello & Sinander, 2025)).

Remark 1. Although we have not yet stated the principal's problem, it is clear that her first-best is the mechanism $(x^0, X^1) \equiv (u^0, u^1)$, which fails to be incentive-compatible due to the conflict of interest ($u^1 < u^0$). If there were no conflict of interest ($u^1 \geq u^0$), then the first-best would be IC.

In the sequel, we equip the set $\mathbf{R}_+$ of times with the Lebesgue measure, so that a 'null set of times' means a set of Lebesgue measure zero, and 'almost everywhere (a.e.)' means 'except possibly on a null set of times'.

Observe that two IC mechanisms (x^0, X^1) and $(x^{0\dagger}, X^1)$ which differ only in that $x^0 \neq x^{0\dagger}$ on a null set are payoff-equivalent.²⁰ For this reason, we shall not distinguish between such mechanisms in the sequel, instead treating them as identical.²¹

3 The principal's problem

In this section, we formulate the principal's problem, and define undominated and optimal mechanisms. We then derive an upper bound on the agent's utility in undominated mechanisms.

3.1 After disclosure

To determine the principal's payoff, we must fill in the gaps in the definition of a mechanism. So fix a mechanism (x^0, X^1) , and suppose that the agent discloses at time t . For each of the remaining periods $s \in [t, \infty)$, the principal must determine

- (1) which technology (F^0 or F^1) will be used, and
- (2) what flow utility $x_s^{1,t}$ the agent will enjoy.

²⁰ x^0 enters payoffs as $\mathbf{E}_G \left(\int_0^\tau e^{-rt} x_t^0 dt \right)$ and $\mathbf{E}_G \left(\int_0^\tau e^{-rt} F^0(x_t^0) dt \right)$, respectively. Modifying x^0 on a null set has no effect on the integrals, and thus leaves both players' payoffs unchanged, no matter what the breakthrough distribution G .

²¹We term such (x^0, X^1) and $(x^{0\dagger}, X^1)$ *versions* of each other. A mechanism is really an equivalence class: a maximal set whose every element is a version of every other.

Part (1) is straightforward: the principal is always (weakly) better off using the new technology.

For (2), the principal must choose a (measurable) utility flow $x^{1,t} : [t, \infty) \rightarrow [0, \infty)$ subject to providing the agent with the continuation utility specified by the mechanism:

$$r \int_t^\infty e^{-r(s-t)} x_s^{1,t} ds = X_t^1.$$

She chooses so as to maximise her post-disclosure payoff

$$r \int_t^\infty e^{-r(s-t)} F^1(x_s^{1,t}) ds.$$

Since the frontier F^1 is concave, the constant flow $x^{1,t} \equiv X_t^1$ is optimal.

Parts (1) and (2) together imply that the principal earns a flow payoff of $F^1(X_t^1)$ forever following a time- t disclosure in a mechanism (x^0, X^1) .

3.2 Undominated and optimal mechanisms

The principal's payoff from an incentive-compatible mechanism (x^0, X^1) is

$$\Pi_G(x^0, X^1) := \mathbf{E}_G \left(r \int_0^\tau e^{-rt} F^0(x_t^0) dt + e^{-r\tau} F^1(X_\tau^1) \right),$$

where the expectation is over the random breakthrough time $\tau \sim G$.²² Her problem is to maximise her payoff by choosing among IC mechanisms.

A basic adequacy criterion for a mechanism is that it not be *dominated* by another mechanism, by which we mean that the alternative mechanism is weakly better under every distribution and strictly better under at least one:

Definition 2. Let (x^0, X^1) and $(x^{0\dagger}, X^{1\dagger})$ be incentive-compatible mechanisms. The former *dominates* the latter iff

$$\Pi_G(x^0, X^1) \geq (>) \Pi_G(x^{0\dagger}, X^{1\dagger}) \quad \text{for every (some) distribution } G.$$

An IC mechanism is *undominated* iff no IC mechanism dominates it.

Domination is a distribution-free concept: the principal weakly prefers a dominating mechanism whatever her belief G about the likely time of the breakthrough. When the principal's belief G makes her exactly indifferent between two mechanisms, one of which dominates the other, choosing the dominating mechanism means maximising the principal's ex-post payoff (which cannot hurt, and seems more prudent if the principal entertains even a little doubt about G).

²²For IC mechanisms (x^0, X^1) such that $X_\tau^1 = \infty$ with positive probability, we interpret ' $F^1(\infty)$ ' as $\lim_{u \uparrow \infty} F^1(u) = -\infty$, so that $\Pi_G(x^0, X^1) := -\infty$.

Definition 3. An incentive-compatible mechanism is *optimal* for a distribution G iff it maximises Π_G and is undominated.

Undominated and optimal mechanisms exist, by standard arguments which we omit (see the working paper (Curello & Sinander, 2025)).

3.3 An upper bound on the agent's utility

Absent incentive concerns, the principal never wishes to give the agent utility strictly exceeding u^0 , since both frontiers are downward-sloping to the right of u^0 . The principal could use utility promises in excess of u^0 as an incentive tool, however. This is never worthwhile:

Lemma 0. Any undominated incentive-compatible mechanism (x^0, X^1) satisfies $x^0 \leq u^0$ almost everywhere.

Proof. Let (x^0, X^1) be an IC mechanism in which $x^0 > u^0$ on a non-null set of times. Consider the alternative mechanism $(\min\{x^0, u^0\}, X^1)$ in which the agent's pre-disclosure flow is capped at u^0 . This mechanism dominates the original one: its pre-disclosure flow is lower, strictly on a non-null set, and the frontier F^0 is strictly decreasing on $[u^0, \infty)$. And it is incentive-compatible: prompt disclosure is as attractive as in the original (IC) mechanism, and disclosing with delay (or never disclosing) is weakly less attractive since the agent earns a lower flow payoff $\min\{x^0, u^0\} \leq x^0$ while delaying. $\square$

4 Keeping the agent indifferent

In this section, we describe how undominated mechanisms incentivise the agent. This result is a stepping stone to the deadline characterisation of undominated mechanisms that we develop in next two sections.

To formulate the agent's problem in a mechanism (x^0, X^1) , let X_t^0 denote the period- t present value of the remainder of the pre-disclosure flow x^0 :

$$X_t^0 := r \int_t^\infty e^{-r(s-t)} x_s^0 ds.$$

In a period t in which the agent has observed but not yet disclosed the breakthrough, she chooses between

- disclosing promptly (payoff X_t^1),
- disclosing with any delay $d > 0$ (payoff $X_t^0 + e^{-rd}(X_{t+d}^1 - X_{t+d}^0)$), and

- never disclosing (payoff X_t^0).

Incentive-compatibility demands precisely that the agent weakly prefer the first option. Our first result asserts that in an undominated mechanism, she must in fact be indifferent between all three alternatives:

Proposition 0. Any undominated incentive-compatible mechanism (x^0, X^1) satisfies $X^0 = X^1$.

That is, the reward X_t^1 for disclosure must equal the present value $X_t^0 = r \int_t^\infty e^{-r(s-t)} x_s^0 ds$ of the remainder of the pre-disclosure flow x^0 .

A naïve intuition for Proposition 0 is that, were the agent strictly to prefer prompt disclosure in some period t , the principal could reduce her disclosure reward X_t^1 without violating IC. The trouble with this idea is that if $X_t^1 \leq u^1$, then lowering X_t^1 would *hurt* the principal (refer to Figure 1 on p. 7). This is no mere quibble, for (as we shall see) undominated mechanisms will spend time in $[0, u^1]$. More broadly, in a general dynamic environment, it is not clear that IC ought to bind everywhere.

The proof is in appendix A. Below, we outline the main idea in discrete time, then highlight the additional details that arise in continuous time.

Sketch proof. Let time $t \in \{0, 1, 2, \dots\}$ be discrete, and write $\beta := e^{-r}$ for the discount factor. A mechanism (x^0, X^1) is incentive-compatible iff in each period s , the agent prefers prompt disclosure to delaying by one period and to never disclosing:

$$\begin{aligned} X_s^1 &\geq (1 - \beta)x_s^0 + \beta X_{s+1}^1 && \text{(delay IC)} \\ X_s^1 &\geq X_s^0 && \text{(non-disclosure IC)} \end{aligned}$$

(Delay IC also deters delay by two or more periods.) We shall show that undominatedness requires that the delay IC inequalities be equalities; we omit the argument that non-disclosure IC must also hold with equality.

So let (x^0, X^1) be an IC mechanism with delay IC slack in some period t :

$$X_t^1 > (1 - \beta)x_t^0 + \beta X_{t+1}^1.$$

Observe that if the terms x_t^0 and X_{t+1}^1 on the right-hand side are $\geq u^1$, then the left-hand side X_t^1 must strictly exceed u^1 . Equivalently, it must be that either

$$(i) \ X_t^1 > u^1, \quad (ii) \ x_t^0 < u^1, \quad \text{or} \quad (iii) \ X_{t+1}^1 < u^1.$$

In each of these cases, we shall find a mechanism that dominates (x^0, X^1) . We will use the fact that non-disclosure IC is slack in each period $s \leq t$.²³

In case (i), the naïve intuition is vindicated: lowering X_t^1 toward u^1 really does improve the principal's payoff (strictly in case of a breakthrough in period t). And this preserves IC: the (slack) period- t delay IC and non-disclosure IC hold for a small enough decrease, while delay IC *slackens* in period $t - 1$ and is unaffected in all other periods, and non-disclosure IC is unaffected in all periods other than t .

In case (ii), increase x_t^0 toward u^1 , by an amount small enough to preserve the (slack) period- t delay IC and period- s non-disclosure IC for each $s \leq t$. Other periods' delay IC is undisturbed, and so is non-disclosure IC in periods $s > t$. Since F^0 increases strictly to the left of $u^1 < u^0$, the principal's payoff improves (strictly in case of a breakthrough after t).

Finally, in case (iii), *increase* X_{t+1}^1 toward u^1 . (The opposite of the naïve intuition.) The principal is better off (strictly in case of a period- $(t + 1)$ breakthrough). Period- t delay IC abides provided the modification is small, while delay IC is loosened in period $t + 1$ and unaffected in other periods. Non-disclosure IC is clearly preserved. $\square$

The proof in appendix A is based on the logic of the sketch above, but must handle two issues that arise in continuous time. First, in case (ii), x^0 must be increased on a *non-null* set of times if the principal's payoff is to increase strictly under some distribution. Secondly, in cases (i) and (iii), it is typically not possible to modify X^1 in a single period while preserving IC.

In light of Proposition 0, an undominated incentive-compatible mechanism (x^0, X^1) is pinned down by the pre-disclosure flow x^0 , since the disclosure reward X^1 must always equal the present value of the remainder of x^0 :

$$X_t^1 = X_t^0 = r \int_t^\infty e^{-r(s-t)} x_s^0 ds \quad \text{for each } t \in \mathbf{R}_+.$$

We therefore drop superscripts in the sequel, writing an IC mechanism simply as (x, X) , where $X_t := r \int_t^\infty e^{-r(s-t)} x_s ds$ for each $t \in \mathbf{R}_+$. Since mechanisms of this form are automatically IC, we refer to them simply as 'mechanisms'. By Lemma 0, we need only consider mechanisms (x, X) that satisfy $x \leq u^0$ a.e.

²³To prove this, use induction on $s \in \{t, t - 1, \dots, 2, 1, 0\}$. In the base case $s = t$, $X_t^1 > (1 - \beta)x_t^0 + \beta X_{t+1}^1 \geq (1 - \beta)x_t^0 + \beta X_{t+1}^0 \equiv X_t^0$ by period- t delay IC (which is slack) and period- $(t+1)$ non-disclosure IC. For the induction step, suppose that period- $(s+1)$ non-disclosure IC is slack, where $s < t$; then $X_s^1 \geq (1 - \beta)x_s^0 + \beta X_{s+1}^1 > (1 - \beta)x_s^0 + \beta X_{s+1}^0 \equiv X_s^0$, where the weak inequality holds by period- s delay IC.

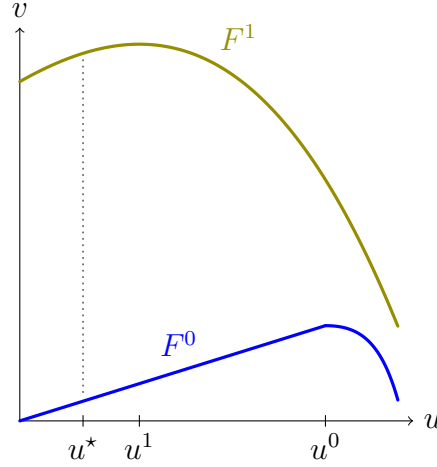


Figure 3: Utility possibility frontiers in the affine case. u^* is where the frontiers are furthest apart.

5 Deadline mechanisms

In this section, we uncover a deadline structure of undominated mechanisms when the old utility possibility frontier F^0 is affine on $[0, u^0]$, as in Figure 3. We further characterise the optimal choice of deadline, given the breakthrough distribution.

We start with the affine case partly for reasons of conceptual clarity: this case lays bare a ‘front-loading’ force that will provide the key to understanding undominated mechanisms in general. The affine case is also important in its own right, since affineness frequently arises in applications, for two basic reasons. The first reason is convexification (recall Figure 2 on p. 8). In the simplest case, with just two allocations, the utility possibility frontier is the straight line connecting the two feasible utility profiles.²⁴ More generally, the utility possibility set is the convex hull of all feasible utility profiles, so its upper boundary F^0 is affine if there are two profiles such that the line segment connecting them lies above all other profiles.

The second reason is that in (utilitarian) policy applications, such as unemployment insurance (§7 below), the agent’s utility directly enters the principal’s payoff in a linear fashion. Explicitly, the agent’s utility is $u = \phi(a)$, where $a \in \mathcal{A}$ is a policy variable and $\phi : \mathcal{A} \rightarrow [0, \infty)$ is surjective, and the principal’s utility is $v = u - \psi(a)$ for some function $\psi : \mathcal{A} \rightarrow \mathbf{R}$, so

$$F^0(u) = \sup_{a \in \mathcal{A}} \{u - \psi(a) : u = \phi(a)\} = u - \inf_{a \in \mathcal{A}} \{\psi(a) : u = \phi(a)\} \quad \text{for each } u \in [0, \infty).$$

²⁴In-between profiles are attained by rapidly switching back and forth (or randomising).

The first term is linear, so if the second term has almost no curvature, then F^0 is approximately affine.²⁵

The utility level u^* (defined in §2) admits a simple description when F^0 is affine: it is the unique $u \in [0, u^0]$ at which the frontiers are furthest apart,²⁶ as indicated in Figure 3. A *deadline mechanism* is one in which the agent's utility absent disclosure is at the efficient level u^0 before a deterministic deadline, and at the inefficiently low level u^* afterwards:

Definition 4. A mechanism (x, X) is a *deadline mechanism* iff

$$x_t = \begin{cases} u^0 & \text{for } t \leq T \\ u^* & \text{for } t > T \end{cases} \quad \text{for some } T \in [0, \infty].$$

Deadline mechanisms are simple: only two utility levels are used, with a single switch between them. And they form a small class of mechanisms, parametrised by a single number: the deadline T . (The utility levels u^0 and u^* are not free parameters, being pinned down by the technologies F^0, F^1 .)

The agent's reward X upon disclosure in a deadline mechanism (equal to the present value of the remainder of the pre-disclosure flow x) is decreasing until the deadline, then constant at u^* :

$$X_t = \begin{cases} (1 - e^{-r(T-t)})u^0 + e^{-r(T-t)}u^* & \text{for } t \leq T \\ u^* & \text{for } t > T. \end{cases} \quad (\diamond)$$

5.1 Only deadline mechanisms are undominated

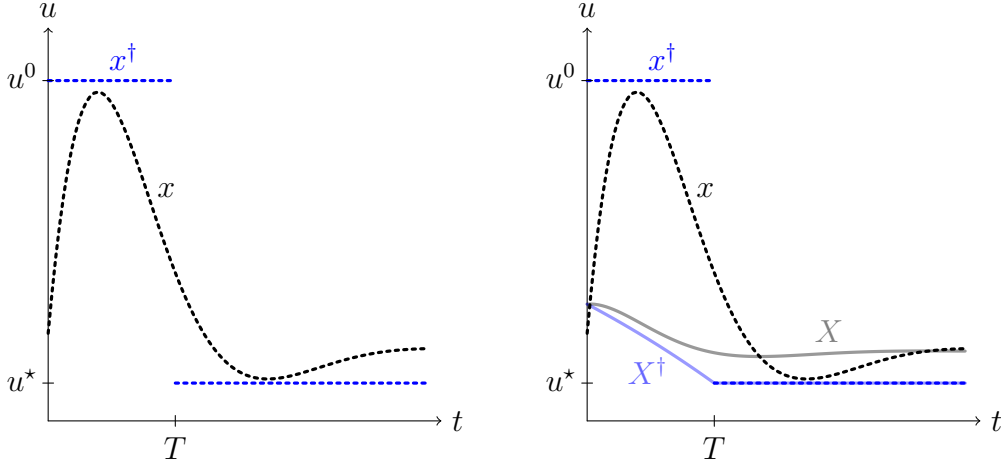
The affine case admits a sharp prediction: no matter what the shapes of the new frontier F^1 and breakthrough distribution G , the principal will choose a mechanism from the small and simple deadline class.

Theorem 1. If the old frontier F^0 is affine on $[0, u^0]$, then any undominated mechanism is a deadline mechanism.

The welfare implications are stark: ex-post Pareto efficiency in case of an early breakthrough, and surplus destruction otherwise. In particular, absent a breakthrough, the old technology is operated Pareto-efficiently (i.e. on the downward-sloping part of F^0 , specifically at u^0) before the deadline, and inefficiently (at u^*) afterwards. Once the new technology arrives, it is deployed

²⁵For example, if $\mathcal{A}$ is a convex subset of $\mathbf{R}$ and ϕ, ψ are twice continuously differentiable with $\phi' > 0 < \psi'$, then $F^0(u) = u - \psi(\phi^{-1}(u))$ for each $u \in [0, \infty)$, so the curvature $|F^{0''}/F^{0'}|$ is small if the curvature difference $|\psi''/\psi' - \phi''/\phi'|$ is small.

²⁶ u^* is a strict local maximum of the gap $F^1 - F^0$, which is concave when F^0 is affine.



(a) $x^\dagger$ is higher early and lower late.

(b) $X^\dagger \leq X$, with equality at 0.

Figure 4: Sketch proof of Theorem 1: front-loading by a deadline mechanism.

efficiently (on the downward-sloping part of F^1) if its arrival was early (while $X \geq u^1$).²⁷ If its arrival was late, then F^1 is operated inefficiently if $u^* < u^1$, and efficiently if $u^* = u^1$. These welfare implications, as well as the special role played by u^* , are general properties that hold even outside of the affine case, so we postpone discussing them fully until §6.2 below.

We prove Theorem 1 in appendix B. Below, we give an intuitive sketch.

Sketch proof. Fix a non-deadline mechanism (x, X) with $x \leq u^0$, and assume for simplicity that $x \geq u^*$. We will show that (x, X) is dominated by the deadline mechanism $(x^\dagger, X^\dagger)$ whose deadline T satisfies

$$\underbrace{(1 - e^{-rT})u^0 + e^{-rT}u^*}_{= X_0^\dagger \text{ by } (\diamond)} = X_0.$$

This mechanism is a *front-loading* of (x, X) : the pre-disclosure flow has the same present value $X_0 = r \int_0^\infty e^{-rt} x_t dt$, but is higher early and lower late, as depicted in Figure 4a. As time passes, the present value $X_t^\dagger = r \int_t^\infty e^{-r(s-t)} x_s^\dagger ds$ of the remainder of the front-loaded flow $x^\dagger$ rapidly diminishes, so that $X^\dagger$ is weakly below X in every period (see Figure 4b).

The principal's period- t continuation payoff if the agent never discloses is

$$Y_t := r \int_t^\infty e^{-r(s-t)} F^0(x_s) ds = F^0 \left(r \int_t^\infty e^{-r(s-t)} x_s ds \right) = F^0(X_t),$$

²⁷A detail: $X_t \geq u^1$ holds in early periods t only if the deadline is sufficiently late. We show in the next section that this must be the case in undominated mechanisms.

where the middle equality holds by the affineness of F^0 . Her payoff is thus

$$\begin{aligned}\Pi_G(x, X) &= \mathbf{E}_G \left(Y_0 - e^{-r\tau} Y_\tau + e^{-r\tau} F^1(X_\tau) \right) \\ &= F^0(X_0) + \mathbf{E}_G \left(e^{-r\tau} [F^1 - F^0](X_\tau) \right).\end{aligned}$$

Front-loading lowers X toward u^* , leaving X_0 unchanged. Since $F^1 - F^0$ is (strictly) decreasing on $[u^*, u^0]$ by definition of u^* , this improves the principal's payoff whatever the distribution G . The improvement is in fact strict for any full-support distribution. Thus $(x^\dagger, X^\dagger)$ dominates (x, X) . $\square$

The key simplification in the above sketch is the assumption that $x \geq u^*$. The proof in appendix B dispenses with this assumption by choosing the deadline T to satisfy $(1 - e^{-rT})u^0 + e^{-rT}u^* = \max\{X_0, u^*\}$, and showing (in a few extra steps) that this yields a dominating mechanism even if $x \not\geq u^*$.

Theorem 1 provides a rationale for deadline mechanisms even when F^0 is not exactly affine: provided F^0 has only moderate curvature, the principal loses little by restricting attention to deadline mechanisms.

5.2 Undominated deadlines

Theorem 1 asserts that only deadline mechanisms are undominated when F^0 is affine, but does not adjudicate between deadlines. In fact, not every deadline mechanism is undominated. Consider a deadline T so early that $X_0 < u^1$. Since the disclosure reward X decreases over time in a deadline mechanism, we have $X_\tau < u^1$ whatever the time τ of the breakthrough.

The principal can do better by using the later deadline $\underline{T}$ that satisfies $X_0 = u^1$, or explicitly (using equation $(\diamond)$ on p. 18)

$$(1 - e^{-r\underline{T}})u^0 + e^{-r\underline{T}}u^* = u^1.$$

This raises the agent's disclosure reward X toward u^1 , improving the principal's post-disclosure payoff $F^1(X_\tau)$ whatever the breakthrough time τ (strictly if $\tau < \underline{T}$). The principal also enjoys the high pre-disclosure flow $F^0(u^0) > F^0(u^*)$ for longer, which is beneficial in case of a late breakthrough.

Undominatedness thus requires a deadline no earlier than $\underline{T}$. This condition is not only necessary, but also sufficient:

Proposition 1. If the old frontier F^0 is affine on $[0, u^0]$, then a mechanism is undominated iff it is a deadline mechanism with deadline $T \in [\underline{T}, \infty]$.

The proof is in appendix C.

5.3 Optimal deadlines

Proposition 1 narrows the search for an optimal mechanism to deadline mechanisms with a sufficiently late deadline. The optimal choice among these depends on the breakthrough distribution G .

A late deadline is beneficial if the breakthrough occurs late, as the efficient high utility u^0 is then provided for a long time. The cost is that in case of an early breakthrough, the agent must be given a utility of $X > u^1$ forever. A first-order condition balances this trade-off:

Proposition 2. Assume that the old frontier F^0 is affine on $[0, u^0]$, that the new frontier F^1 is differentiable on $(0, u^0)$ with bounded derivative, and that $u^* > 0$. A mechanism is optimal for G iff it is a deadline mechanism and satisfies $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$.

In other words, the new technology should be operated optimally *on average*. This is a restriction on the deadline T because X is a function of it, as described by equation $(\diamond)$ on p. 18. Indeed, it implies comparative statics: optimal deadlines become later when the breakthrough distribution G becomes later in the sense of first-order stochastic dominance.²⁸ This improves the agent's ex-ante payoff X_0 , as can be seen from equation $(\diamond)$.

Proposition 2 is proved in appendix G.

6 Optimal mechanisms in general

In this section, we show that optimal mechanisms in the general (non-affine) case exhibit a graduated deadline structure: absent disclosure, the agent's utility still declines from u^0 toward u^* , but not necessarily abruptly. Given the breakthrough distribution, we describe the optimal path.

To shorten proofs, we shall impose a well-behavedness assumption. The results remain true if this assumption is dropped: see the working paper (Curello & Sinander, 2025).

Definition 5. We say that the model primitives F^0, F^1, G are *well-behaved* iff F^0 and F^1 are differentiable on $(0, u^0)$ with bounded derivatives, and either (i) F^0 is strictly concave on $[0, u^0]$ or (ii) F^1 is strictly concave on $[0, u^0]$ and G has full support.

²⁸To see why, recall that F^1 is concave, and observe (from $(\diamond)$) that a deadline mechanism's disclosure reward X is decreasing over time and increases pointwise when the deadline becomes later. This comparative-statics result remains true if F^1 is not differentiable; see the working paper (Curello & Sinander, 2025).

6.1 Qualitative features of optimal mechanisms

Recall from §2 that u^* denotes the greatest $u \in [0, u^0]$ at which the old and new frontiers F^0, F^1 have equal slopes, as depicted in Figure 1 (p. 7).

Theorem 2. Let G be a distribution with $G(0) = 0$ and unbounded support. Assume that F^0, F^1, G are well-behaved. Then any mechanism (x, X) that is optimal for G has x decreasing

$$\text{from } \lim_{t \rightarrow 0} x_t = u^0 \quad \text{toward} \quad \lim_{t \rightarrow \infty} x_t = u^*.$$
²⁹

That is, optimal mechanisms are just like deadline mechanisms, except that the transition from u^0 to u^* may be gradual. This graduality follows directly from relaxing affineness: when F^0 has a strictly concave shape, by definition, the principal prefers providing intermediate utility to providing only the extreme utilities u^*, u^0 . Theorem 2 is the combination of this mechanical effect with the front-loading insight expressed by Theorem 1.

The proof is in appendix E. As mentioned above, Theorem 2 remains true if the well-behavedness assumption is dropped, at the cost of a longer proof; see the working paper (Curello & Sinander, 2025).

The role of monotonicity is *not* to provide incentives: on the contrary, mechanisms of the form (x, X) satisfy IC (with equality) by definition, whatever the pre-disclosure flow $x : \mathbf{R}_+ \rightarrow [0, u^0]$. Rather, what Theorem 2 asserts is that if x is not decreasing, then there is a better mechanism. This claim is non-trivial to prove.

Absent a breakthrough, efficiency deteriorates as we travel leftward along the upward-sloping part of the old frontier F^0 . Once the new technology becomes available, it is operated efficiently (on the downward-sloping part of F^1) if its arrival was sufficiently early.³⁰ If its arrival was late, then F^1 is operated inefficiently if $u^* < u^1$, and efficiently if $u^* = u^1$.

The distributional hypotheses are mild: $G(0) = 0$ means that the new technology is unavailable initially, while unbounded support rules out an effectively finite horizon. The former's role is as a sufficient condition for $\lim_{t \rightarrow 0} x_t = u^0$, while the latter is required by our proof strategy.

²⁹Recall that a mechanism has multiple *versions* (footnote 21, p. 12). Theorem 2 asserts that any optimal mechanism has a version with the stated properties. We focus on $\lim_{t \rightarrow 0} x_t$ rather than x_0 because ' $x_0 = u^0$ ' is vacuous: *any* mechanism has a version satisfying it.

³⁰We show in appendix F that $X_t > u^1$ holds in all sufficiently early periods t .

6.2 Discussion

Two salient features of Theorems 1 and 2 are the special role played by u^* and the possibility (in case of a late breakthrough) of perpetual surplus destruction. We now discuss these two properties.

For simplicity, assume that F^0 and F^1 are differentiable, and consider a mechanism that is eventually constant: $x = \bar{u}$ on (T, ∞) , where $\bar{u} \in (0, u^0)$ and $G(T) < 1$. Unless $\bar{u} = u^*$, the mechanism (x, X) may be improved by a simple perturbation:

$$x^\varepsilon = \begin{cases} x & \text{on } [0, T] \\ \bar{u} + \varepsilon & \text{on } (T, T + \ln(2)/r] \\ \bar{u} - \varepsilon & \text{on } [T + \ln(2)/r, \infty) \end{cases} \quad \text{where } \varepsilon \neq 0.$$

If $\varepsilon > 0$, then this is a ‘front-loading’, making the pre-disclosure flow x higher early on (before $T + \ln(2)/r$) and lower later, while keeping $X^\varepsilon = X$ on $[0, T]$.³¹ Since $\frac{d}{d\varepsilon}x^\varepsilon|_{\varepsilon=0} = \frac{d}{d\varepsilon}X^\varepsilon|_{\varepsilon=0} = 0$ on $[0, T]$, perturbing ε away from zero changes the principal’s payoff $\Pi_G(x^\varepsilon, X^\varepsilon)$ at rate

$$\begin{aligned} & \left. \frac{d}{d\varepsilon} \mathbf{E}_G \left(r \int_0^\tau e^{-rt} F^0(x_t^\varepsilon) dt \right) \right|_{\varepsilon=0} + \left. \frac{d}{d\varepsilon} \mathbf{E}_G \left(e^{-r\tau} F^1(X_\tau^\varepsilon) \right) \right|_{\varepsilon=0} \\ &= \mathbf{E}_G \left(r \int_0^\tau e^{-rt} \left. \frac{d}{d\varepsilon} x_t^\varepsilon \right|_{\varepsilon=0} dt \right) \times F^{0'}(\bar{u}) + K_G \times F^{1'}(\bar{u}) \\ &= K_G \times [F^{1'}(\bar{u}) - F^{0'}(\bar{u})] \quad \text{where } K_G := \mathbf{E}_G \left(e^{-r\tau} \left. \frac{d}{d\varepsilon} X_\tau^\varepsilon \right|_{\varepsilon=0} \right), \end{aligned}$$

where the second equality holds since the big expectation equals $\mathbf{E}_G(\phi'_\tau(0))$ where $\phi_\tau(\varepsilon) := r \int_0^\tau e^{-rt} x_t^\varepsilon dt = X_0^\varepsilon - e^{-r\tau} X_\tau^\varepsilon$. Thus *whatever* the breakthrough distribution G , the principal’s payoff can be improved by perturbing ε except if $F^{0'}(\bar{u}) = F^{1'}(\bar{u})$, or equivalently $\bar{u} = u^*$.

This accounts for the special role of u^* . It also implies the optimality of perpetual surplus destruction in case of a late breakthrough (after T), since setting $\bar{u} = u^* < u^1$ yields $X = x < u^1$ on (T, ∞) .

Economically, the above argument boils down to a demonstration that u^* balances the cost and benefit of ‘front-loading’, so that neither front-loading ($\varepsilon > 0$) nor ‘back-loading’ ($\varepsilon < 0$) yields an improvement. The benefit of front-loading is that the pre-disclosure flow x is experienced only before the breakthrough, so making it higher early and lower late is mechanically better.³² The cost of front-loading is that it lowers the disclosure reward X , thereby increasing the severity of perpetual surplus destruction in case of a late breakthrough.

³¹Because $X_T^\varepsilon = X_T + \varepsilon e^{rT} (\int_T^{T+\ln(2)/r} r e^{-rs} ds - \int_{T+\ln(2)/r}^\infty r e^{-rs} ds) = X_T$ for each ε .

³²The principal prefers a higher pre-disclosure flow since F^0 is increasing on $[0, u^0]$.

6.3 Optimal transition

Theorem 2 describes the distribution-free qualitative features of optimal mechanisms, but does not specify the precise manner in which the agent's utility ought to decline from u^0 toward u^* . The optimal path, for a given breakthrough distribution, is characterised by an Euler equation:

Proposition 3. Let G be a distribution with $G(0) = 0$ and unbounded support. Assume that $u^* > 0$, and that F^0, F^1, G are well-behaved. Then any mechanism (x, X) that is optimal for G satisfies the initial condition $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$ and the Euler equation

$$F^{0'}(x_t) \geq \mathbf{E}_G(F^{1'}(X_\tau) | \tau > t) \quad \text{for each } t \in \mathbf{R}_+, \text{ with equality if } x_t < u^0.^{33}$$

The initial condition $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$ demands that the new technology be used optimally on average, just like the first-order condition for an optimal deadline in the affine case (Proposition 2, p. 21). The special role of u^* can be deduced from the Euler equation: as $t \rightarrow \infty$, $X_t = r \int_t^\infty e^{-r(s-t)} x_s ds$ converges to $\underline{u} := \lim_{t \rightarrow \infty} x_t$, so $F^{0'}(\underline{u}) = F^{1'}(\underline{u})$, which is to say that $\underline{u} = u^*$.

The proof is in appendix F. Proposition 3 remains true if well-behavedness is weakened to differentiability of F^0, F^1 on $(0, u^0)$; see the working paper (Curello & Sinander, 2025).

Sketch proof. A mechanism (x, X) with $0 < x < u^0$ may be perturbed near an arbitrary period $t \in \mathbf{R}_+$ by adding ε to x on $[t, t + \delta)$, where $\varepsilon \neq 0$ and $\delta > 0$ are small. This changes $X_s = r \int_s^\infty e^{-r(s'-s)} x_{s'} ds'$ for $s \leq t$ by $re^{-r(t-s)}\delta\varepsilon + o(\delta\varepsilon)$, so changes the principal's payoff $\Pi_G(x, X)$ by

$$re^{-rt}F^{0'}(x_t)\delta\varepsilon[1 - G(t)] + \int_{[0,t]} e^{-rs}F^{1'}(X_s) \left(re^{-r(t-s)}\delta\varepsilon\right) G(ds) + o(\delta\varepsilon).$$

If (x, X) is optimal, then it cannot be improved by such perturbations:

$$F^{0'}(x_t)[1 - G(t)] + \int_{[0,t]} F^{1'}(X_s)G(ds) = 0. \quad (\mathcal{E}_t)$$

Letting $t \rightarrow \infty$ yields $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$. Substituting this equality into $(\mathcal{E}_t)$ and dividing by $1 - G(t) > 0$ yields $F^{0'}(x_t) = \mathbf{E}_G(F^{1'}(X_\tau) | \tau > t)$. $\square$

³³Here $F^{j'}(0)$ ($F^{j'}(u^0)$) for $j \in \{0, 1\}$ denotes the right-hand (left-hand) derivative. Recall that a mechanism has multiple *versions* (footnote 21, p. 12). In full, the proposition asserts that some (any) version satisfies the Euler equation for (almost) every $t \in \mathbf{R}_+$.

To understand the Euler equation, differentiate it and rearrange to obtain

$$\dot{x}_t = - \underbrace{\left(\frac{G'(t)}{1 - G(t)} \right)}_{\text{hazard rate}} \underbrace{\frac{F^{0'}(x_t) - F^{1'}(X_t)}{-F^{0''}(x_t)}}_{\text{curvature}}.^{34}$$

Thus the agent's pre-disclosure utility declines in proportion to the hazard rate, and in inverse proportion to the local curvature of the old frontier F^0 . As the latter would suggest, outside of the well-behaved case, x jumps over any affine segments ($F^{0''} = 0$ and ' $\dot{x} = \infty$ '), and pauses at kinks ($F^{0''} = -\infty$ and $\dot{x} = 0$).

As for comparative statics, it can be shown (see the working paper (Curello & Sinander, 2025)) that as the breakthrough distribution G becomes later in the sense of monotone likelihood ratio, the disclosure reward X increases in every period. (The pre-disclosure flow x need not increase pointwise.) It follows in particular that the agent's ex-ante payoff X_0 improves.

Although our focus is on general properties, there are special cases in which the Euler equation may be solved in closed form:

Example 2. Let the breakthrough arrive at constant rate $\lambda > 0$, so that $G(t) = 1 - e^{-\lambda t}$ for every $t \in \mathbf{R}_+$. Fix $u^1 < u^0$ in $(0, \infty)$, and assume that

$$F^j(u) := a^j \left(u^j - \frac{1}{2}u \right) u + b^j \quad \text{for each } j \in \{0, 1\} \text{ and every } u \in [0, u^0],$$

where $0 < a^0 < a^1 > a^0 u^0 / u^1$, and $b^1 - b^0$ is large enough that $F^1 \geq F^0$. Solving the Euler equation yields the optimal mechanism x given by

$$x_t := (u^0 - u^*) e^{-\lambda k t} + u^* \quad \text{for each } t \in \mathbf{R}_+,$$

where

$$u^* = \frac{a^1 u^1 - a^0 u^0}{a^1 - a^0} \quad \text{and} \quad k := \frac{1 + r/\lambda}{2} \left(\sqrt{\frac{r/\lambda}{\left(\frac{1+r/\lambda}{2}\right)^2} \left(\frac{a^1}{a^0} - 1 \right) + 1} - 1 \right).$$

In the special case $r = \lambda$, this simplifies to $k = \sqrt{a^1/a^0} - 1$.

7 Application to unemployment insurance

If unemployment benefits are generous but time-limited, then a worker who receives a job offer before her benefits run out may have an incentive to delay

³⁴This expression is valid under the additional assumptions that G admits a continuous density and that F^0 possesses a continuous and strictly negative second derivative.

starting her new job, for example by arranging a deferred start date (or by simply waiting before accepting the offer). Empirically, such strategic delay appears to be widespread.³⁵

In this section, we study the design of unemployment insurance (‘UI’) schemes when workers can exercise such strategic delay. We focus in particular on the merits of *deadline benefit schemes*, in which the short-term unemployed receive a generous benefit, while those remaining unemployed past a deadline see their benefit reduced to a much lower level. Such schemes are used in many countries, including Germany, France and Sweden. We also study the optimal choice of deadline.

Related literature. The literature on optimal unemployment insurance has two main strands. The first concerns the moral-hazard problem of incentivising job-search effort (Shavell & Weiss, 1979; Hopenhayn & Nicolini, 1997). We contribute to the second strand, which studies the adverse-selection problem arising from privately observed job offers (Atkeson & Lucas, 1995).³⁶ (It can be shown that our conclusions in this section would not change if we added moral hazard to the model: see the working paper (Curello & Sinander, 2025).) Within this second strand, our contribution is to characterise optimal UI under the assumption that workers can delay starting a new job, rather than having to start right away.

7.1 Model

A worker (agent) is unemployed. At a random time $\tau \sim G$, she receives a job offer. If she accepts, then she chooses when to start. The worker’s ability to delay her start date is the distinguishing feature of our otherwise-standard model. The state observes in real time whether the worker is employed, but cannot observe whether she has received a job offer. All jobs are permanent and pay the same wage, denoted $w > 0$.

The worker’s utility is $u = \phi(C) - \kappa(L)$, where $C \geq 0$ is her consumption and $L \geq 0$ her labour supply. We assume that $\phi, \kappa : [0, \infty) \rightarrow [0, \infty)$ are respectively strictly concave and strictly convex, that they are differentiable on $(0, \infty)$ with strictly positive derivatives that satisfy

$$\lim_{C \rightarrow \infty} \phi'(C) = 0, \quad \lim_{C \rightarrow 0} \phi'(C) = \infty \quad \text{and} \quad \lim_{L \rightarrow 0} \kappa'(L) = 0,$$

³⁵See Boone and van Ours (2012), DellaVigna, Lindner, Reizer, and Schmieder (2017), and Kyyrä, Pesola, and Verho (2019).

³⁶See also Thomas and Worrall (1990), Atkeson and Lucas (1992), Hansen and İmrohoroglu (1992), and Shimer and Werning (2008).

and that they are continuous with $\phi(0) = \kappa(0) = 0$ and $\lim_{C \rightarrow \infty} \phi(C) = \infty$. We interpret $C = 0$ as the lowest socially acceptable standard of living. If the worker is unemployed, then $L = 0$.

The state controls unemployment benefits and income taxes. Following the literature, we impose no constraints on policy:³⁷ income taxation after re-employment can be non-linear, for example, and can depend on the length of the preceding unemployment spell. These policy instruments can implement any allocation (C, L) which the worker prefers to autarky.³⁸ We may therefore model the state as directly choosing (C, L) , subject to $u = \phi(C) - \kappa(L) \geq 0$.

The state's objective is social welfare $v = u + \lambda \times (wL - C)$, where u is the worker's welfare, $wL - C$ is net tax revenue, and $\lambda > 0$ is the shadow value of public funds. The utility possibility frontiers for unemployed and employed workers are thus

$$F^0(u) := \max_{C \geq 0} \{u + \lambda(-C) : \phi(C) = u\}$$

$$\text{and } F^1(u) := \max_{C, L \geq 0} \{u + \lambda(wL - C) : \phi(C) - \kappa(L) = u\},$$

respectively. These frontiers satisfy our model assumptions (§2):

Lemma 1. In the application to unemployment insurance, the frontiers F^0, F^1 are strictly concave and continuous, with unique peaks u^0, u^1 that satisfy $u^1 < u^0$. The gap $F^1 - F^0$ is strictly decreasing, so that $u^* = 0$.

The conflict of interest $u^1 < u^0$ arises because the social first-best requires employed workers to supply labour ($L > 0$), which they dislike, without compensating them with extra consumption (first-best consumption is $C^0 := (\phi')^{-1}(\lambda)$ regardless of employment status). This is an instance of the fact, well-known in public finance since Mirrlees (1971, 1974),³⁹ that welfare-maximisation (absent incentive constraints) does not ‘reward merit’: on the contrary, it dictates efficient production, meaning that more productive workers work harder. The proof of Lemma 1 is elementary but tedious, so we omit it.

We shall use the term ‘unemployment insurance (UI) scheme’ for a mechanism. By Proposition 0 (p. 15), undominated schemes keep the worker only just willing promptly to start a job, so have the form (x, X) . Implicit in a

³⁷This has been the standard approach since Hopenhayn and Nicolini (1997).

³⁸An unemployed worker's consumption is simply her benefit. To get an employed worker to choose a bundle (C, L) satisfying $u := \phi(C) - \kappa(L) \geq 0$, use the income tax schedule $\theta(Y) = \min\{Y, mY + b\}$, with $m, b \in \mathbf{R}$ chosen so that the worker's income $L' \mapsto wL' - \theta(wL')$ is tangent at L to her indifference curve $L' \mapsto \phi^{-1}(\kappa(L') + u)$.

³⁹See the third section of Mirrlees (1974), as well as p. 201 of Mirrlees (1971).

UI scheme (x, X) are the benefit B_t paid to the time- t unemployed (given by $x_t = \phi(B_t)$) and the labour supply L_t and tax bill $\theta_t = wL_t - C_t$ of a worker who started working at t (which satisfy $X_t = \phi(wL_t - \theta_t) - \kappa(L_t)$).

7.2 Optimal unemployment insurance

Optimal UI schemes are described by Theorem 2 (p. 22): unemployment benefits $B_t = \phi^{-1}(x_t)$ decrease over time, from $C^0 = \phi^{-1}(u^0)$ toward $0 = \phi^{-1}(u^*)$. Thus workers enjoy socially optimal consumption at the beginning of an unemployment spell, but see their benefits reduced over time, with the long-term unemployed provided only with society's lowest acceptable standard of living ('zero consumption').

Employed workers are rewarded with a higher continuation utility X_t the earlier they start a job. This involves a mix of lower labour supply and more generous tax treatment of earnings (yielding higher consumption).

A *deadline UI scheme* is one in which a generous benefit of C^0 is paid to the short-term unemployed, while those remaining unemployed beyond a deadline receive a low benefit just sufficient to finance the minimum standard of living ('zero consumption'). Such schemes are widespread in practice, used in e.g. Germany, France and Sweden.

Our results speak to the desirability of such deadline schemes. Theorem 1 (p. 18) implies that a deadline scheme is approximately optimal if F^0 is close to affine, a condition which is satisfied if the worker's consumption utility ϕ has limited curvature or if the social value λ of tax revenue is moderate. Conversely, if neither assumption is close to being satisfied, then our results predict substantial welfare gains from more gradual tapering, as in Italy.

Given the prevalence of deadline schemes (whatever their merits), the choice of deadline is an important policy problem. Our analysis highlights labour-market prospects as a key consideration: a worker with worse chances (a later job-finding distribution G , in the sense of first-order stochastic dominance) should be set a later deadline.⁴⁰ Two implications are that older workers ought to face later deadlines and that extensions should be granted during recessions. These recommendations are broadly followed in Germany and France, where workers older than about 50 face more lenient deadlines, and all workers' deadlines were prolonged during the 2020 recession.

⁴⁰In particular, the optimal deadline described by Proposition 2 (p. 21) is later when G is, as noted at the end of §5.3.

Appendices

A Proof of Proposition 0 (p. 15)

We shall follow the sketch proof, but with significant elaborations aimed at overcoming the two technical hurdles discussed at the end of §4.

For any mechanism (x^0, X^1) , let $h : \mathbf{R}_+ \rightarrow [-\infty, \infty]$ be given by $h(t) := e^{-rt}(X_t^1 - X_t^0)$ for each $t \in \mathbf{R}_+$.⁴¹ Proposition 0 asserts precisely that undominated IC mechanisms have h identically equal to zero.

Observation 1. A mechanism (x^0, X^1) is incentive-compatible exactly if h is (a) decreasing and (b) non-negative.

Proof. Part (a) (part (b)) of the definition of incentive-compatibility on p. 11 requires precisely that h be decreasing (non-negative). $\square$

Continuity lemma. Any undominated IC mechanism has h continuous.

Proof. We prove the contrapositive. Fix an IC mechanism (x^0, X^1) .

Suppose that h is discontinuous at some $t \in (0, \infty)$. Since h is decreasing and X^0 is continuous, $\lim_{s \uparrow t} X_s^1$ and $\lim_{s \downarrow t} X_s^1$ exist and satisfy $\lim_{s \uparrow t} X_s^1 \geq X_t^1 \geq \lim_{s \downarrow t} X_s^1$, with one of the inequalities strict. We shall assume that

$$\lim_{s \uparrow t} X_s^1 = X_t^1 > \lim_{s \downarrow t} X_s^1,$$

omitting the similar arguments for the other two cases. If $\lim_{s \downarrow t} X_s^1 < u^1$, then we may increase X^1 toward u^1 on a small interval $(t, t + \varepsilon)$ while keeping h decreasing.⁴² If instead $\lim_{s \downarrow t} X_s^1 \geq u^1$, then $\lim_{s \uparrow t} X_s^1 = X_t^1 > u^1$, so that we may decrease X^1 toward u^1 on a small interval $(t - \varepsilon, t]$ while keeping h decreasing.⁴³ In either case, IC is preserved, and the principal's payoff Π_G is (strictly) increased under any (full-support) distribution G .

Suppose instead that h is discontinuous at $t = 0$; then $X_0^1 > \lim_{s \downarrow 0} X_s^1$ by IC and the continuity of X^0 . The case $\lim_{s \downarrow 0} X_s^1 < u^1$ may be dealt with as above. If $\lim_{s \downarrow 0} X_s^1 \geq u^1$, then lowering X_0^1 toward $\lim_{s \downarrow 0} X_s^1$ preserves IC and (strictly) increases Π_G for any distribution G (with $G(0) > 0$). $\square$

⁴¹In case $X_t^1 = X_t^0 = \infty$, we let $h(t) := 0$ by convention.

⁴²Choose an $\varepsilon > 0$ small enough that $X^1 + \varepsilon < \min\{u^1, X_t^1\}$ on $(t, t + \varepsilon)$. Let $X_s^{1\uparrow} := X_s^1 - (s - t) + \varepsilon$ for $s \in (t, t + \varepsilon)$ and $X^{1\uparrow} := X^1$ off $(t, t + \varepsilon)$. Then $X^1 \leq X^{1\uparrow} \leq u^1$, with the first inequality strict on $(t, t + \varepsilon)$. We have $h^\uparrow \geq h \geq 0$, and $h^\uparrow$ is clearly decreasing on $[0, t]$ and on (t, ∞) . At t , we have $h^\uparrow(t) - \lim_{s \downarrow t} h^\uparrow(s) = e^{-rt}(X_t^1 - \lim_{s \downarrow t} X_s^1 - \varepsilon) \geq 0$.

⁴³Choose an $\varepsilon \in (0, 1/r)$ small enough that $X^1 - \varepsilon > \lim_{s \downarrow t} X_s^1$ and $h > \varepsilon$ on $(t - \varepsilon, t]$. Let $X_s^{1\downarrow} := X_s^1 + t - s - \varepsilon$ for $s \in (t - \varepsilon, t]$ and $X^{1\downarrow} := X^1$ off $(t - \varepsilon, t]$. Then $u^1 \leq X^{1\downarrow} \leq X^1$, with the second inequality strict on $(t - \varepsilon, t]$. Clearly $h^\downarrow$ is non-negative, and is decreasing on $[0, t - \varepsilon]$ and on (t, ∞) . It is decreasing on $[t - \varepsilon, t]$ since $h^\downarrow(s) - h(s) = e^{-rs}(t - s - \varepsilon)$ is (by our choice of $\varepsilon < 1/r$). And at t , $h^\downarrow(t) - \lim_{s \downarrow t} h^\downarrow(s) = e^{-rt}(X_t^1 - \varepsilon - \lim_{s \downarrow t} X_s^1) \geq 0$.

Proof of Proposition 0. Let (x^0, X^1) be an IC mechanism, so that h is non-negative and decreasing, and suppose that h is not identically zero. By the continuity lemma, we may assume that h (and thus X^1) is continuous.

We consider three cases. (The first two concern slack ‘delay IC’: Case 1 [Case 2] corresponds to the sketch proof’s case (ii) [cases (i) and (iii)]. Case 3 is where ‘delay IC’ binds, but ‘non-disclosure IC’ is slack.) In each case, we shall construct an incentive-compatible mechanism $(x^{0\ddagger}, X^{1\ddagger})$ such that

$$\Pi_G(x^{0\ddagger}, X^{1\ddagger}) \geq (>) \Pi_G(x^0, X^1) \quad \text{for every (full-support) } G. \quad (\text{D})$$

Define $A := \{t \in \mathbf{R}_+ : h \text{ is differentiable at } t \text{ and } h'(t) < 0\}$.

Case 1: $\{t \in A : x_t^0 < u^0\}$ is non-null. Since $h > 0$ on A ,⁴⁴ there is an $\varepsilon > 0$ for which the set

$$A_\varepsilon := \{t \in A : x_t^0 + \varepsilon < u^0, h(t) \geq \varepsilon \text{ and } h'(t) + r\varepsilon \leq 0\}$$

is non-null.⁴⁵ Define $x^{0\ddagger} := x^0 + \varepsilon \mathbf{1}_{A_\varepsilon}$, and consider the mechanism $(x^{0\ddagger}, X^1)$. Clearly $x^0 \leq x^{0\ddagger} \leq u^0$, and $x^{0\ddagger} \neq x^0$ on the non-null set A_ε , so that (D) holds by the strict monotonicity of F^0 on $[0, u^0]$. $h^\ddagger$ is decreasing since for any $t < t'$ in $\mathbf{R}_+$,

$$\begin{aligned} h^\ddagger(t') - h^\ddagger(t) &= h(t') - h(t) + r\varepsilon \int_t^{t'} e^{-rs} \mathbf{1}_{A_\varepsilon}(s) ds \\ &\leq \int_t^{t'} h' \mathbf{1}_{A_\varepsilon} + r\varepsilon \int_t^{t'} e^{-rs} \mathbf{1}_{A_\varepsilon}(s) ds \leq 0, \end{aligned}$$

where the first inequality holds since h is decreasing,⁴⁶ and the second holds by definition of A_ε . As for non-negativity, we have $h^\ddagger = h \geq 0$ on $(\sup A_\varepsilon, \infty)$, while $h^\ddagger \geq 0$ on $[0, \sup A_\varepsilon)$ since $h^\ddagger$ is decreasing and $h^\ddagger \geq h - \varepsilon \geq 0$ on A_ε by definition of the latter. Thus $(x^{0\ddagger}, X^1)$ is incentive-compatible.

Case 2: There are $t' < t''$ in $\mathbf{R}_+$ such that $h(t') > h(t'')$ and $X^1 \neq u^1$ on $[t', t'']$. Since X^1 is continuous, we have either $X^1 > u^1$ on $[t', t'']$ or $X^1 < u^1$ on $[t', t'']$. We shall assume the former, omitting the similar argument for the latter case. Because $s \mapsto e^{rs}h(t'') + X_s^0$ is continuous and takes the value $X_{t''}^1 > u^1$ at $s = t''$,

$$t^* := \inf \left\{ t \in [t', t''] : e^{rs}h(t'') + X_s^0 \geq u^1 \text{ for all } s \in [t, t''] \right\}$$

⁴⁴Since $h \geq 0$, $h(t) = 0$ implies $\liminf_{t' \downarrow t} [h(t') - h(t)] / (t' - t) \geq 0$ and thus $t \notin A$.

⁴⁵ $A_0 = \bigcup_{n \in \mathbf{N}} A_{1/n}$ is non-null, so continuity of measures (with λ denoting the Lebesgue measure) yields $0 < \lambda(A_0) = \lim_{n \rightarrow \infty} \lambda(A_{1/n})$, whence $\lambda(A_{1/n}) > 0$ for some $n \in \mathbf{N}$.

⁴⁶Recall the Lebesgue decomposition $h = h_a + h_s$ where h_a is decreasing and absolutely continuous and h_s is decreasing with $h'_s = 0$ a.e. (e.g. Stein & Shakarchi, 2005, p. 150).

is well-defined and strictly smaller than t'' . Define

$$X_t^{1\dagger} := \begin{cases} e^{rt}h(t'') + X_t^0 & \text{for } t \in [t^*, t'') \\ X_t^1 & \text{for } t \notin [t^*, t''), \end{cases}$$

and consider the mechanism $(x^0, X^{1\dagger})$. This mechanism is IC since $h^\dagger = h + [h(t'') - h]\mathbf{1}_{[t^*, t'')}$ is clearly decreasing and non-negative.

It remains to show that $(x^0, X^{1\dagger})$ satisfies (D). Since X^1 and $X^{1\dagger}$ differ only on $[t^*, t'')$ and F^1 is strictly decreasing on $[u^1, \infty)$, it suffices to prove that

$$u^1 \leq X_t^{1\dagger} \leq (<) X_t^1 \quad \text{for every (some) } t \in [t^*, t'').^{47}$$

The first inequality holds by definition of t^* . For the second, observe that

$$X_t^{1\dagger} - X_t^1 = e^{rt} [h^\dagger(t) - h(t)] = e^{rt} [h(t'') - h(t)] \leq 0 \quad \text{for } t \in [t^*, t'')$$

since h is decreasing. We claim that the inequality is strict at $t = t^*$. If $t^* = t'$, then this is true because $h(t') > h(t'')$. And if not, then $t^* \in (t', t'')$, in which case $X_{t^*}^{1\dagger} = u^1 < X_{t^*}^1$ by continuity of X^0 and $X^1 > u^1$.

Case 3: neither Case 1 nor Case 2. Since X^1 is continuous, every $t \in \mathbf{R}_+$ belongs either to a maximal open interval on which $X^1 \neq u^1$ or else to a maximal closed interval on which $X^1 = u^1$. h is increasing on any interval of the former kind since we are not in Case 2. We shall show that h is also increasing on each interval of the latter kind; then since h is continuous, it is increasing and thus constant.

So fix an interval I of the latter kind. Since h is decreasing, its derivative $h'(t) = re^{-rt}(x_t^0 - u^1)$ exists a.e. on I . As we are not in Case 1, we have for a.e. $t \in I$ that either $h'(t) = 0$ or $x_t^0 = u^0$, and in the latter case $h'(t) = re^{-rt}(u^0 - u^1) > 0$. Assuming wlog that $x^0 \leq u^0$,⁴⁸ the expression for h' implies that h is ru^0 -Lipschitz on I . Thus h is increasing on I , as desired.

Since (by hypothesis) h is not identically zero, it is constant at some $k > 0$, so that $X_t^1 = X_t^0 + e^{rt}k$ for every $t \in \mathbf{R}_+$. Thus $X^{1\dagger} := \min\{X^1, X^0 + u^1\}$ is strictly smaller than X^1 after some time $T > 0$, so that $(x^0, X^{1\dagger})$ satisfies (D). And it is incentive-compatible.⁴⁹ $\square$

⁴⁷It is enough for the inequality to be strict at a single time $t \in [t^*, t'')$, since it then holds strictly on a proper interval by the continuity of X^1 and $X^{1\dagger}$ on $[t^*, t'')$.

⁴⁸Otherwise the IC mechanism $(\min\{x^0, u^0\}, X^1)$ would satisfy (D).

⁴⁹We have $h^\dagger(t) = e^{-rt}u^1 \in (0, h^\dagger(T))$ for $t > T$, and this expression is decreasing.

B Proof of Theorem 1 (p. 18)

Fix a non-deadline mechanism (x, X) with $x \leq u^0$ a.e.;⁵⁰ we will show that it is dominated by the deadline mechanism $(x^\dagger, X^\dagger)$ whose deadline T satisfies

$$(1 - e^{-rT})u^0 + e^{-rT}u^\star \equiv X_0^\dagger = X_0 \vee u^\star,$$

where ‘ $\vee$ ’ denotes the pointwise maximum.

Claim. $X^\dagger \leq X \vee u^\star$.

Proof. For $t \geq T$, we have $X^\dagger = u^\star \leq X \vee u^\star$. For $t < T$, suppose first that $X_0^\dagger = X_0$; then since $x^\dagger = u^0 \geq x$ on $[0, t] \subseteq [0, T]$, we have

$$e^{-rt}X_t^\dagger = X_0^\dagger - r \int_0^t e^{-rs}x_s^\dagger ds \leq X_0 - r \int_0^t e^{-rs}x_s ds = e^{-rt}X_t \leq e^{-rt}(X_t \vee u^\star).$$

If instead $X_0^\dagger = u^\star$, then the fact that $x^\dagger \geq u^\star$ yields

$$e^{-rt}X_t^\dagger = X_0^\dagger - r \int_0^t e^{-rs}x_s^\dagger ds \leq u^\star - r \int_0^t e^{-rs}u^\star ds = e^{-rt}u^\star \leq e^{-rt}(X_t \vee u^\star).$$

□

The concave function $F^1 - F^0$ is uniquely maximised at $u^\star$, so is strictly increasing on $[0, u^\star]$ and strictly decreasing on $[u^\star, u^0]$. Since $u^\star \leq X^\dagger \leq X \vee u^\star$ by the claim, it follows that

$$[F^1 - F^0](X^\dagger) \geq [F^1 - F^0](X \vee u^\star). \quad (1)$$

Since $X \vee u^\star \geq X$, and the two differ only when both are in $[0, u^\star]$, we have

$$[F^1 - F^0](X \vee u^\star) \geq [F^1 - F^0](X), \quad (2)$$

which chained together with the preceding inequality yields

$$[F^1 - F^0](X^\dagger) \geq [F^1 - F^0](X). \quad (3)$$

The facts that $X_0^\dagger = X_0 \vee u^\star \geq X_0$ and that F^0 is increasing on $[0, u^0]$ together imply

$$F^0(X_0^\dagger) \geq F^0(X_0). \quad (4)$$

⁵⁰IC mechanisms not of this form are dominated, by Lemma 0 and Proposition 0.

Thus for any distribution G , using the expression for the principal's payoff derived in the sketch proof (p. 20), we have

$$\begin{aligned}
\Pi_G(x^\dagger, X^\dagger) &= F^0(X_0^\dagger) + \mathbf{E}_G(e^{-r\tau} [F^1 - F^0](X_\tau^\dagger)) \\
&\geq F^0(X_0^\dagger) + \mathbf{E}_G(e^{-r\tau} [F^1 - F^0](X_\tau)) && \text{by (3)} \\
&\geq F^0(X_0) + \mathbf{E}_G(e^{-r\tau} [F^1 - F^0](X_\tau)) && \text{by (4)} \\
&= \Pi_G(x, X).
\end{aligned}$$

It remains show that $(x^\dagger, X^\dagger)$ delivers a *strict* improvement for some distribution G . We shall accomplish this by showing that the inequality (3) holds strictly on a non-null set of times, so that the first inequality in the above display is strict for any distribution G with full support. Since $X^\dagger \leq X \vee u^*$ by the claim and $X, X^\dagger$ are continuous, there are two cases: either (a) $X^\dagger < X \vee u^*$ on a non-null set of times, or (b) $X^\dagger = X \vee u^*$.

Case (a): $X^\dagger < X \vee u^*$ on a non-null set $\mathcal{T}$. In this case, the inequality (1) holds strictly on $\mathcal{T}$, and thus so does (3).

Case (b): $X^\dagger = X \vee u^*$. Since the original mechanism (x, X) is not a deadline mechanism, there must be a non-null set of times on which $x \neq x^\dagger$, and thus $X \neq X^\dagger = X \vee u^*$ on some non-null set $\mathcal{T}$, so that $X < X \vee u^*$ on $\mathcal{T}$. Then (2) is strict on $\mathcal{T}$, and thus so is (3). $\square$

C Proof of Proposition 1 (p. 20)

Write (x^T, X^T) for the deadline mechanism with deadline T , and $\pi_G(T)$ for its payoff under a distribution G . By Theorem 1, any undominated mechanism is a deadline mechanism. We showed in the text (§5.2, p. 20) that those with deadline $T < \underline{T}$ are dominated, so it remains only to show that those with deadline $T \geq \underline{T}$ are not. We shall rely on the following claim.

Claim. If the deadline mechanism (x^T, X^T) is dominated for some $T \geq \underline{T}$, then it is dominated by another deadline mechanism.

Proof of the claim. Fix a $T \geq \underline{T}$ such that (x^T, X^T) is dominated by some IC mechanism (x^0, X^1) ; we must show that (x^T, X^T) is dominated by a deadline mechanism. We may assume without loss that $x^0 \leq u^0$, since if $x^0 > u^0$ on a non-null set of times, then we may replace (x^0, X^1) with the IC mechanism $(x^{0\dagger}, X^{1\dagger})$ obtained from the proof of Lemma 0, which satisfies $x^{0\dagger} \leq u^0$ and dominates (x^0, X^1) , hence dominates (x^T, X^T) .

The proof of Theorem 1 (appendix B) shows that any IC mechanism (x^0, X^1) that satisfies $x^0 \leq u^0$ a.e. and $X^0 = X^1$ is either a deadline mechanism or is dominated by a deadline mechanism. It therefore suffices to show that $X^0 = X^1$.

For each $t \in \mathbf{R}_+$, let G^t denote the point mass at t . Note that

$$F^1(X_0^T) = \pi_{G^0}(T) \leq \Pi_{G^0}(x^0, X^1) = F^1(X_0^1),$$

where the inequality holds since (x^T, X^T) is dominated by (x^0, X^1) . Then $X_0^1 \leq X_0^T$ since $X_0^T \geq u^1$ (as $T \geq \underline{T}$) and F^1 is concave with unique peak u^1 . Moreover,

$$\begin{aligned} F^0(X_0^T) &= \int_0^\infty r e^{-rt} F^0(x_t^T) dt = \lim_{t \rightarrow \infty} \pi_{G^t}(T) \\ &\leq \limsup_{t \rightarrow \infty} \Pi_{G^t}(x^0, X^1) \leq \int_0^\infty r e^{-rt} F^0(x_t^0) dt = F^0(X_0^0), \end{aligned}$$

where the first and last equalities hold since F^0 is affine, the second equality since F^1 is bounded on $[0, u^0]$ and $X^T \leq u^0$, the first inequality since (x^T, X^T) is dominated by (x^0, X^1) , and the second inequality since F^1 is bounded above. Then $X_0^T \leq X_0^0$ since $X_0^T \leq u^0$ and F^0 is strictly increasing on $[0, u^0]$. Altogether, we have shown that $X_0^1 \leq X_0^T \leq X_0^0$. Since (x^0, X^1) is IC, it follows by Observation 1 in appendix A (p. 29) that $X^0 = X^1$. $\square$

By the claim, it suffices to prove that (x^T, X^T) for $T \in [\underline{T}, \infty]$ is not dominated by another deadline mechanism.

Part 1: finite deadlines. Fix a deadline $T \in [\underline{T}, \infty)$; we shall identify a distribution G under which the deadline T yields a strictly higher payoff than any other deadline. In particular, consider the point mass at $T - \underline{T}$. The mechanism (x^T, X^T) has $x = u^0$ on $[0, T - \underline{T}] \subseteq [0, T]$ and $X_{T-\underline{T}}^T = (1 - e^{-r\underline{T}})u^0 + e^{-r\underline{T}}u^* = u^1$ by $(\diamond)$ on p. 18 and the definition of $\underline{T}$. Thus (x^T, X^T) provides flow payoff $F^0(u^0)$ before the breakthrough and $F^1(u^1)$ afterwards, which is the first-best. Any other deadline T' has $X_{T'-\underline{T}}^{T'} \neq u^1$, so provides a strictly lower post-disclosure payoff and a no higher pre-disclosure payoff.

Part 2: the infinite deadline. Fix an arbitrary finite deadline $T \in [0, \infty)$; we must show that (x^T, X^T) does not dominate (x^∞, X^∞) . To that end, we shall identify a distribution G under which the former mechanism is strictly worse. In particular, let G^t denote the point mass at some $t \geq T$. Under this

distribution, the payoff difference between the two mechanisms is

$$\begin{aligned} \pi_{G^t}(T) - \pi_{G^t}(\infty) &= e^{-rt} \left\{ \left[F^1(u^*) - F^1(u^0) \right] - \left[F^0(u^*) - F^0(u^0) \right] \right\} \\ &\quad + e^{-rT} \left[F^0(u^*) - F^0(u^0) \right]. \end{aligned}$$

The second term is strictly negative since F^0 is uniquely maximised at u^0 and $u^* \leq u^1 < u^0$. By choosing $t \geq T$ large enough, we can make the first term as small as we wish, so that the payoff difference is strictly negative. $\square$

D An Euler equation

In this appendix, we argue that optimal mechanisms are described by an Euler equation, and that this equation admits a decreasing solution. These results will be used in the next two appendices to prove Theorem 2 and Proposition 3 (pp. 22 and 24).

We assume throughout this appendix that F^0 and F^1 are differentiable on $(0, u^0)$ with bounded derivatives $F^{0'}$ and $F^{1'}$. Extend both continuously to $[0, u^0]$.

D.1 The Euler equation and optimality

Let $\mathcal{X}$ be the set of all measurable maps $\mathbf{R}_+ \rightarrow [0, u^0]$.

Definition 6. Given a distribution G , a mechanism (x, X) with $x \in \mathcal{X}$ satisfies the Euler equation (for G) iff for a.e. $t \in \mathbf{R}_+$,

$$[1 - G(t)] F^{0'}(x_t) + \int_{[0, t]} F^{1'}(X_s) G(ds) \leq (\geq) 0 \quad \text{if } x_t < u^0 \text{ (if } x_t > 0). \quad (\text{E})$$

For a given breakthrough distribution G , define $\pi_G : \mathcal{X} \rightarrow \mathbf{R}$ by

$$\pi_G(x) := \Pi_G(x, X) = \mathbf{E}_G \left(r \int_0^\tau e^{-rs} F^0(x_s) ds + e^{-r\tau} F^1(X_\tau) \right).$$

This is the principal's payoff under G from the mechanism (x, X) .

Euler lemma. For a mechanism (x, X) with $x \in \mathcal{X}$ and a distribution G , x belongs to $\arg \max_{\mathcal{X}} \pi_G$ iff (x, X) satisfies the Euler equation for G .

Proof. Note first that for all $x, x' \in \mathcal{X}$, the Gateaux derivative of π_G at x in

the direction $x' - x$ is

$$\begin{aligned}
D\pi_G(x, x' - x) &= \lim_{\alpha \downarrow 0} \frac{\pi_G(x + \alpha[x' - x]) - \pi_G(x)}{\alpha} \\
&= \mathbf{E}_G \left(r \int_0^\tau e^{-rt} F^{0'}(x_t) [x'_t - x_t] dt + e^{-r\tau} F^{1'}(X_\tau) [X'_\tau - X_\tau] \right) \\
&= \mathbf{E}_G \left(r \int_0^\infty e^{-rt} [\mathbf{1}_{(t, \infty)}(\tau) F^{0'}(x_t) + \mathbf{1}_{[0, t]}(\tau) F^{1'}(X_\tau)] [x'_t - x_t] dt \right) \\
&= r \int_0^\infty e^{-rt} \left[[1 - G(t)] F^{0'}(x_t) + \int_{[0, t]} F^{1'}(X_s) G(ds) \right] [x'_t - x_t] dt,
\end{aligned}$$

where the third equality follows by the bounded convergence theorem since $F^{0'}$ and $F^{1'}$ are bounded.

For the ‘if’ part, suppose that (x, X) satisfies the Euler equation. Then $D\pi_G(x, x' - x) \leq 0$ for any $x' \in \mathcal{X}$. Since F^0, F^1 are concave and the map $x \mapsto X$ is linear, π_G is concave. Thus for any $\alpha \in (0, 1)$ and $x' \in \mathcal{X}$, we have

$$\pi_G(x') - \pi_G(x) \leq \frac{\pi_G(x + \alpha[x' - x]) - \pi_G(x)}{\alpha},$$

so that letting $\alpha \downarrow 0$ yields $\pi_G(x') - \pi_G(x) \leq D\pi_G(x, x' - x) \leq 0$. So $x \in \arg \max_{\mathcal{X}} \pi_G$.

For the ‘only if’ part, we prove the contra-positive: suppose that $x \in \mathcal{X}$ does not satisfy the Euler equation; we will show that $x \notin \arg \max_{\mathcal{X}} \pi_G$. It suffices to exhibit an $x' \in \mathcal{X}$ such that $D\pi_G(x, x' - x) > 0$. If there exists a non-null $A \subseteq \mathbf{R}_+$ on which (E) fails and $x < u^0$ holds, then choose $x' := x + (u^0 - x) \mathbf{1}_A$. If not, then there exists a non-null $A \subseteq \mathbf{R}_+$ on which (E) fails and $x > 0$ holds; in this case, choose $x' := x \mathbf{1}_{\mathbf{R}_+ \setminus A}$. $\square$

For $x > 0$, the backward-looking integral equation (E) is equivalent to a forward-looking integral equation plus an initial condition:

Lemma 2. For any $x \in \mathcal{X}$ with $x > 0$, (x, X) satisfies the Euler equation iff $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$ and, for a.e. $t \in \mathbf{R}_+$ with $G(t) < 1$,

$$F^{0'}(x_t) \geq \mathbf{E}_G(F^{1'}(X_\tau) | \tau > t) \quad \text{with equality if } x_t < u^0. \quad (5)$$

Proof. Let $\psi(t) := [1 - G(t)] F^{0'}(x_t) + \int_{[0, t]} F^{1'}(X_s) G(ds)$ for each $t \in \mathbf{R}_+$. For any $t \in \mathbf{R}_+$, $\int_{(t, \infty)} F^{1'}(X_s) G(ds)$ is finite since $F^{1'}$ is bounded, so we may add and subtract it to obtain

$$\psi(t) = \begin{cases} [1 - G(t)] [F^{0'}(x_t) - \mathbf{E}_G(F^{1'}(X_\tau) | \tau > t)] + \mathbf{E}_G(F^{1'}(X_\tau)) & \text{if } G(t) < 1 \\ \mathbf{E}_G(F^{1'}(X_\tau)) & \text{if } G(t) = 1. \end{cases}$$

Thus if $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$ holds and (5) holds for a.e. $t \in \mathbf{R}_+$ with $G(t) < 1$, then (x, X) satisfies the Euler equation. For the converse, suppose that (x, X) satisfies the Euler equation; we will show that $\lim_{t \rightarrow \infty} \text{ess inf}_{s \geq t} \psi(s) = 0$. This is sufficient since it implies that there is a sequence $(t^n)_{n \in \mathbf{N}}$ in $\mathbf{R}_+$ along which $t^n \rightarrow \infty$ and $\psi(t^n) \rightarrow 0$ as $n \rightarrow \infty$, so that $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$ by bounded convergence, as $F^{0'}$ and $F^{1'}$ are bounded, which implies that (5) holds for a.e. $t \in \mathbf{R}_+$ with $G(t) < 1$.

Since (x, X) satisfies the Euler equation and $x > 0$, we have $\psi \geq 0$ a.e. and $\psi(t) = 0$ for a.e. $t \in \mathbf{R}_+$ such that $x_t < u^0$. It follows immediately that if there is no $T' \in \mathbf{R}_+$ such that $x = u^0$ a.e. on $[T', \infty)$, then $\lim_{t \rightarrow \infty} \text{ess inf}_{s \geq t} \psi(s) = 0$. Assume for the remainder that there exists a $T' \in \mathbf{R}_+$ such that $x = u^0$ a.e. on $[T', \infty)$, and let T be the smallest such T' . It suffices to show that $\psi(t) \leq 0$ holds for every $t > T$ such that $x_t = u^0$.

Note that $T > 0$, because otherwise $X = u^0$, which since $F^{1'}(u^0) < 0$ would imply that (E) fails for sufficiently large $t \in \mathbf{R}_+$. Choose an increasing sequence $(t^n)_{n \in \mathbf{N}}$ in $\mathbf{R}_+$ converging to T along which (E) and $x < u^0$ both hold. Then for all $t > T$ with $x_t = u^0$,

$$\psi(t) \leq [1 - G(T)]F^{0'}(u^0) + \int_{[0, T]} F^{1'}(X_s)G(ds) \leq \limsup_{n \rightarrow \infty} \psi(t^n) \leq 0,$$

where the first inequality holds since $F^{1'}(X_s) = F^{1'}(u^0) \leq 0 \leq F^{0'}(u^0)$ for all $s \geq T$, and the second inequality holds since $F^{0'}(x_{t^n}) \geq F^{0'}(u^0)$ for each $n \in \mathbf{N}$ (as $x^{t^n} < u^0$, and $F^{0'}$ is decreasing) and $F^{1'}(X_T) = F^{1'}(u^0) \leq 0$. $\square$

D.2 Constructing a solution of the Euler equation

Let $\mathcal{X}'$ be the set all of decreasing maps $\mathbf{R}_+ \rightarrow [u^*, u^0]$, endowed with the topology of pointwise convergence.

Existence corollary. For any distribution G with unbounded support, there is a mechanism (x, X) with $x \in \mathcal{X}'$ which satisfies the Euler equation for G .

We prove the existence corollary in two steps. We first show that it holds for a particular class of technologies F^0, F^1 and distributions G (Lemma 3 below), then extend the claim via a series of limit arguments.

Say that F^0, F^1 are *simple* iff they possess bounded derivatives on $(0, u^0)$, $F^{0'}$ is strictly decreasing with Lipschitz continuous inverse, $F^{1'}$ is Lipschitz continuous, and $u^* > 0$.

Lemma 3. Suppose that F^0 and F^1 are simple. Then for any distribution G such that the support $\text{supp } G$ is bounded and G has an atom at $\max \text{supp } G$, there is a mechanism (x, X) with $x \in \mathcal{X}'$ which satisfies the Euler equation for G .

We first prove the existence corollary using Lemma 3, and then prove Lemma 3. Both proofs use the following simple fact.

Observation 2. If a sequence $(x^n)_{n \in \mathbf{N}}$ in $\mathcal{X}$ converges pointwise to $x \in \mathcal{X}$ as $n \rightarrow \infty$, then $F^{0'}(x^n) \rightarrow F^{0'}(x)$ and $F^{1'}(X^n) \rightarrow F^{1'}(X)$ pointwise as $n \rightarrow \infty$.

Proof of Observation 2. Since F^0, F^1 are concave, their derivatives $F^{0'}, F^{1'}$ are continuous. By the bounded convergence theorem, $X^n \rightarrow X$ pointwise as $n \rightarrow \infty$. $\square$

Proof of the existence corollary. Consider two cases.

Case 1: F^0, F^1 are simple. For any $n \in \mathbf{N}$, and let $G^n := \mathbf{1}_{[0,n)}G + \mathbf{1}_{[n,\infty)}$, and observe that since G has unbounded support by hypothesis, Lemma 3 delivers an $x^n \in \mathcal{X}'$ such that (x^n, X^n) satisfies the Euler equation for G^n . By the Helly selection theorem (e.g. Rudin, 1976, p. 167), we may assume (passing to a subsequence if necessary) that $(x^n)_{n \in \mathbf{N}}$ converges to some $x \in \mathcal{X}'$. Then, for any $t \in \mathbf{R}_+$ such that $t < n$,

$$\begin{aligned} & [1 - G^n(t)] F^{0'}(x_t^n) + \int_{[0,t]} F^{1'}(X_s^n) G^n(ds) \\ &= [1 - G(t)] F^{0'}(x_t^n) + \int_{[0,t]} F^{1'}(X_s^n) G(ds) \\ &\rightarrow [1 - G(t)] F^{0'}(x_t) + \int_{[0,t]} F^{1'}(X_s) G(ds) \quad \text{as } n \rightarrow \infty, \end{aligned}$$

where convergence follows from Observation 2 and the bounded convergence theorem. Since (x^n, X^n) satisfies the Euler equation for G^n for each $n \in \mathbf{N}$ and $x^n \rightarrow x$ pointwise as $n \rightarrow \infty$, it follows that (x, X) satisfies the Euler equation for G .

Case 2: F^0, F^1 are arbitrary. Choose a sequence $(F_n^0, F_n^1)_{n \in \mathbf{N}}$ of technologies satisfying the following:

- (a) for each $n \in \mathbf{N}$, F_n^0, F_n^1 are simple, $u_n^* \geq u^*$, and $u_n^0 = u^0$,
- (b) $F_n^{0'} \geq F^{0'}$ and $F_n^{1'} \geq F^{1'}$ for all $n \in \mathbf{N}$,
- (c) $(F_n^{1'})_{n \in \mathbf{N}}$ is uniformly bounded, and
- (d) for all $u \in (0, u^0)$, $F_n^{0'} \rightarrow F^{0'}$ and $F_n^{1'} \rightarrow F^{1'}$ uniformly on $[u, u^0]$.

By (a) and Case 1 above, there exists for each $n \in \mathbf{N}$ an $x^n \in \mathcal{X}'$ such that (x^n, X^n) satisfies the Euler equation for (F_n^0, F_n^1, G) . By the Helly selection theorem (e.g. Rudin, 1976, p. 167), we may assume (passing to a subsequence

if necessary) that $(x^n)_{n \in \mathbf{N}}$ converges pointwise to some $x \in \mathcal{X}'$. For each $t \in \mathbf{R}_+$, define

$$\begin{aligned}\psi^n(t) &:= [1 - G(t)] F_n^{0'}(x_t^n) + \int_{[0,t]} F_n^{1'}(X_s^n) G(ds) \quad \text{for each } n \in \mathbf{N}, \text{ and} \\ \psi(t) &:= [1 - G(t)] F^{0'}(x_t) + \int_{[0,t]} F^{1'}(X_s) G(ds).\end{aligned}$$

Since (x^n, X^n) satisfies the Euler equation for (F_n^0, F_n^1, G) for every $n \in \mathbf{N}$, it suffices to show that $\liminf_{n \rightarrow \infty} \psi^n(t) \geq \psi(t)$ for all $t \in \mathbf{R}_+$ and that $\lim_{n \rightarrow \infty} \psi^n(t) = \psi(t)$ for all $t \in \mathbf{R}_+$ such that $x_t > 0$, since then (x, X) satisfies the Euler equation for (F^0, F^1, G) .

For the former, fix a $t \in \mathbf{R}_+$. Since $F_n^{0'}(x_t^n) \geq F^{0'}(x_t^n)$ for all $n \in \mathbf{N}$ by (b), $\liminf_{n \rightarrow \infty} F_n^{0'}(x_t^n) \geq F^{0'}(x_t)$ as $F^{0'}$ is continuous. Similarly, $\liminf_{n \rightarrow \infty} F_n^{1'}(X_s^n) \geq F^{1'}(X_s)$ for every $s \in [0, t]$. Hence $\liminf_{n \rightarrow \infty} \psi^n(t) \geq \psi(t)$ by Fatou's lemma, which is applicable by (c).

For the latter, fix a $t \in \mathbf{R}_+$ such that $x_t > 0$. Choose a $u \in (0, x_t)$, further choose an $N \in \mathbf{N}$ large enough that $x_t^n \geq u$ for all $n \geq N$, and note that

$$\left| F_n^{0'}(x_t^n) - F^{0'}(x_t) \right| \leq \sup_{u' \in [u, u^0]} \left| F_n^{0'}(u') - F^{0'}(u') \right| + \left| F^{0'}(x_t^n) - F^{0'}(x_t) \right| \quad \text{for all } n \geq N.$$

Letting $n \rightarrow \infty$ yields $F_n^{0'}(x_t^n) \rightarrow F^{0'}(x_t)$, by (d) and Observation 2. Similarly, since $X > 0$ on $[0, t]$ (because x is decreasing), $F_n^{1'}(X_s^n) \rightarrow F^{1'}(X_s)$ as $n \rightarrow \infty$ for each $s \in [0, t]$. Then $\lim_{n \rightarrow \infty} \psi^n(t) = \psi(t)$ by the bounded convergence theorem, which is applicable by (c). $\square$

Proof of Lemma 3. Let $T := \max \text{supp } G \in \mathbf{R}_+$. We shall prove that for each $\alpha \in [u^*, u^0]$, there exists a unique $x^\alpha \in \mathcal{X}'$ satisfying (5) (p. 36) for each $t < T$ and $x_t^\alpha = \alpha$ for each $t \geq T$. Taking this claim for granted for the time being, define $\psi : [u^*, u^0] \rightarrow \mathbf{R}$ by

$$\psi(\alpha) := \mathbf{E}_G \left(F^{1'}(X_\tau^\alpha) \right) \quad \text{for each } \alpha \in [u^*, u^0].$$

It suffices to show that there is an $\alpha \in [u^*, u^0]$ such that $\psi(\alpha) = 0$, since then (x^α, X^α) satisfies the Euler equation for G by Lemma 2 and the fact that $x^\alpha \geq u^* > 0$.

Note that the constant map $t \mapsto u^0$ satisfies (5) for all $t < T$, since $F^{0'}(u^0) \geq 0 \geq F^{1'}(u^0)$. Thus x^{u^0} is constant at u^0 . Similarly, the constant map $t \mapsto u^*$ satisfies (5) for all $t < T$, since $F^{0'}(u^*) = F^{1'}(u^*)$ as $u^* > 0$. Hence x^{u^*} is constant at u^* . Therefore

$$\psi(u^*) = F^{1'}(u^*) = F^{0'}(u^*) \geq 0 \geq F^{1'}(u^0) = \psi(u^0).$$

Note also that for any convergent sequence $(\alpha_n)_{n \in \mathbf{N}}$ in $[u^\star, u^0]$ along which $(x^{\alpha_n})_{n \in \mathbf{N}}$ converges in $\mathcal{X}'$ to some $x \in \mathcal{X}'$, we have $x = x^{\lim_{n \rightarrow \infty} \alpha_n}$ since x satisfies (5) for all $t < T$, by Observation 2 and the bounded convergence theorem. Since $\mathcal{X}'$ is sequentially compact, it follows that the map $[u^\star, u^0] \rightarrow \mathcal{X}'$ given by $\alpha \mapsto x^\alpha$ is continuous,⁵¹ so ψ is likewise continuous, by Observation 2 and the bounded convergence theorem. Hence by the intermediate value theorem, there is an $\alpha \in [u^\star, u^0]$ such that $\psi(x^\alpha) = 0$, as desired.

It remains to prove the existence and uniqueness of x^α for each $\alpha \in [u^\star, u^0]$. To this end, fix an $\alpha \in [u^\star, u^0]$, and let $\mathcal{X}'_\alpha$ be the set of all $x \in \mathcal{X}'$ such that $x = \alpha$ on $[T, \infty)$. Extend the inverse

$$\text{inv } F^{0'} : [F^{0'}(u^0), F^{0'}(0)] \rightarrow [0, u^0]$$

of $F^{0'}$ to $\mathbf{R}$ by letting $\text{inv } F^{0'}$ be constant on $(-\infty, F^{0'}(u^0)]$ and on $[F^{0'}(0), \infty)$. Given any $x \in \mathcal{X}'_\alpha$, let $\mathcal{H}x : \mathbf{R}_+ \rightarrow [0, \infty)$ be given by

$$(\mathcal{H}x)_t := \begin{cases} \text{inv } F^{0'}(\mathbf{E}_G(F^{1'}(X_\tau)|_{\tau > t})) & \text{if } t < T \\ \alpha & \text{if } t \geq T, \end{cases}$$

and note that $\mathcal{H}x \in \mathcal{X}'_\alpha$ since $\mathcal{H}x$ is decreasing as $\text{inv } F^{0'}$ and $F^{1'}$ are, bounded above by u^0 since $\text{inv } F^{0'}$ is, and bounded below by $u^\star$ since $\text{inv } F^{0'}$ is decreasing and

$$\mathbf{E}_G(F^{1'}(X_\tau)|_{\tau > t}) \leq F^{1'}(u^\star) = F^{0'}(u^\star), \quad (6)$$

where the inequality holds since $F^{1'}$ is decreasing, and the equality holds since $u^\star > 0$.

Observe that for any $x \in \mathcal{X}'_\alpha$ and $t < T$, (5) implies $x_t = (\mathcal{H}x)_t$.⁵² Conversely, for any $x \in \mathcal{X}'_\alpha$ and $t < T$, $x_t = (\mathcal{H}x)_t$ implies (5) since

$$\mathbf{E}_G(F^{1'}(X_\tau)|_{\tau > t}) \leq F^{0'}(u^\star) \leq F^{0'}(0)$$

by (6) and the fact that $F^{0'}$ is decreasing.⁵³ It thus suffices to show that the map $\mathcal{H} : \mathcal{X}'_\alpha \rightarrow \mathcal{X}'_\alpha$ has exactly one fixed point.

⁵¹If $\alpha \mapsto x^\alpha$ were not continuous, then we could choose a sequence $(\alpha_n)_{n \in \mathbf{N}}$ in $[u^\star, u^0]$ converging to some $\alpha \in [u^\star, u^0]$ along which $(x^{\alpha_n})_{n \in \mathbf{N}}$ does not converge to x^α . Then for some $t \in \mathbf{R}_+$, $x_t^{\alpha_n}$ converges along a subsequence to some $u \in [u^\star, u^0] \setminus \{x_t^\alpha\}$. By the sequential compactness of $\mathcal{X}'$, there are sub-subsequences along which $(x^{\alpha_n})_{n \in \mathbf{N}}$ converges pointwise, but none has limit x^α .

⁵²If $x_t = u^0$, then $F^{0'}(u^0) \geq \mathbf{E}_G(F^{1'}(X_\tau)|_{\tau > t})$ by (5), so $(\mathcal{H}x)_t = \text{inv } F^{0'}(F^{0'}(u^0)) = u^0 = x_t$, where the first equality holds since $\text{inv } F^{0'}$ is constant on $(-\infty, F^{0'}(u^0)]$.

⁵³This is immediate if $F^{0'}(u^0) \leq \mathbf{E}_G(F^{1'}(X_\tau)|_{\tau > t}) \leq F^{0'}(0)$. If $\mathbf{E}_G(F^{1'}(X_\tau)|_{\tau > t}) < F^{0'}(u^0)$, then $x_t = (\mathcal{H}x)_t$ implies $x_t = u^0$ as $\text{inv } F^{0'}$ is constant on $(-\infty, F^{0'}(u^0)]$, so (5) holds.

Since F^0, F^1 are simple, we may choose an $\ell > 0$ such that $\text{inv } F^{0'}$ and $F^{1'}$ are ℓ -Lipschitz. Let

$$G(T-) := \begin{cases} \lim_{t \uparrow T} G(t) & \text{if } T > 0 \\ 0 & \text{if } T = 0 \end{cases}$$

and $k := \ell^2/[1 - G(T-)]$. Define $\rho : \mathcal{X}'_\alpha \times \mathcal{X}'_\alpha \rightarrow \mathbf{R}_+$ by

$$\rho(x, x^*) = \sup_{t \in [0, T]} e^{kG(t)} |x_t - x_t^*| \quad \text{for all } x, x^* \in \mathcal{X}'_\alpha.$$

It is easy to see that ρ is a metric on $\mathcal{X}'_\alpha$.

Claim. $\mathcal{H}$ is a contraction on the metric space $(\mathcal{X}'_\alpha, \rho)$.

Since ρ is equivalent to the supremum metric, the metric space $(\mathcal{X}'_\alpha, \rho)$ is complete. Thus by the claim and the Banach fixed-point theorem, $\mathcal{H}$ has exactly one fixed point.

Proof of the claim. The result is immediate if $T = 0$, so assume that $T > 0$. Define $\text{inv } G(z) := \min\{t \in \mathbf{R}_+ : G(t) \geq z\}$ for each $z \in [0, 1]$, and note that for all $x, x^* \in \mathcal{X}'_\alpha$ and $z \in [0, 1]$,

$$\begin{aligned} e^{kz} |X_{\text{inv } G(z)} - X_{\text{inv } G(z)}^*| &\leq e^{kz} \sup_{t \in [\text{inv } G(z), T]} |x_t - x_t^*| \\ &\leq e^{kG(\text{inv } G(z))} \sup_{t \in [\text{inv } G(z), T]} |x_t - x_t^*| \leq \rho(x, x^*), \end{aligned}$$

where the last inequality holds since G is increasing. Thus for all $x, x^* \in \mathcal{X}'_\alpha$,

$$\begin{aligned} \rho(\mathcal{H}x, \mathcal{H}x^*) &\leq \sup_{t \in [0, T)} ke^{kG(t)} \int_{(t, \infty)} |X - X^*| dG \\ &= \sup_{t \in [0, T)} ke^{kG(t)} \int_{G(t)}^1 e^{-kz} e^{kz} |X_{\text{inv } G(z)} - X_{\text{inv } G(z)}^*| dz \\ &\leq \left(\sup_{t \in [0, T)} ke^{kG(t)} \int_{G(t)}^1 e^{-kz} dz \right) \rho(x, x^*) = (1 - e^{-k[1-G(0)]}) \rho(x, x^*), \end{aligned}$$

where the first inequality holds since $\text{inv } F^{0'}$ and $F^{1'}$ are ℓ -Lipschitz and $(\mathcal{H}x)_T = (\mathcal{H}x^*)_T$. Since $0 < e^{-k[1-G(0)]} < 1$, this shows that $\mathcal{H}$ is a contraction. $\square$

With the claim established, the proof is complete. $\square$

E Proof of Theorem 2 (p. 22)

We shall argue as follows. Fix an optimal mechanism (x, X) . We first show that if x is decreasing, then $\lim_{t \rightarrow 0} x_t = u^0$ and $\lim_{t \rightarrow \infty} x_t = u^*$ (Lemma 4 below). We then show that x is indeed decreasing, using the Euler equation (appendix D, p. 35), which (x, X) must satisfy by the Euler lemma (appendix D, p. 35).

Recall from appendix D that $F^{0'}$ and $F^{1'}$ denote the derivatives of F^0 and F^1 on $(0, u^0)$, extended continuously to $[0, u^0]$. Also recall from appendix D.2 the definition of $\mathcal{X}'$.

Lemma 4. Suppose that F^0 and F^1 are differentiable on $(0, u^0)$ with bounded derivatives. Let (x, X) with $x \in \mathcal{X}'$ satisfy the Euler equation for some G with unbounded support and $G(0) = 0$. Then $\lim_{t \rightarrow 0} x_t = u^0$ and $\lim_{t \rightarrow \infty} x_t = u^*$.

Proof. Since x is decreasing with $u^* \leq x \leq u^0$, the limits

$$\bar{u} := \lim_{t \rightarrow 0} x_t \quad \text{and} \quad \underline{u} := \lim_{t \rightarrow \infty} x_t$$

exist and satisfy $u^* \leq \underline{u} \leq \bar{u} \leq u^0$.

To show that $\bar{u} \geq u^0$, assume toward a contradiction that $\bar{u} < u^0$. Then $x < u^0$ since x is decreasing. Since $G(0) = 0$, letting $t \rightarrow 0$ in (E) (p. 35) then yields $F^{0'}(\bar{u}) \leq 0$, which is impossible since F^0 is concave and strictly increasing on $[0, u^0]$.

To show that $\underline{u} \leq u^*$, note first that this is immediate if $\underline{u} = 0$. Assume for the remainder that $\underline{u} > 0$, so that $x > 0$ since x is decreasing. Then Lemma 2 (p. 36) yields $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$ and

$$F^{0'}(x_t) = \mathbf{E}_G(F^{1'}(X_\tau) | \tau > t) \leq F^{1'}(\underline{u}) \quad (7)$$

for a.e. $t \in \mathbf{R}_+$ such that $x_t < u^0$, where the equality holds since G has unbounded support and the inequality holds since F^1 is concave and $X \geq \underline{u}$. Note that $x_t < u^0$ for all sufficiently large t , since otherwise $X = u^0$ as x is decreasing, which would contradict $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$. Hence letting $t \rightarrow \infty$ in (7) yields $F^{0'}(\underline{u}) \leq F^{1'}(\underline{u})$, which implies that $\underline{u} \leq u^*$ by definition of the latter. $\square$

Recall from appendix D the definitions of $\mathcal{X}$ and π_G , the Euler lemma, and the existence corollary.

Proof of Theorem 2. Let G be a distribution with unbounded support and $G(0) = 0$, and assume that F^0, F^1, G are well-behaved; we must show that (x, X) has the properties listed in Theorem 2.

By the existence corollary, there is a mechanism $(x^\dagger, X^\dagger)$ with $x^\dagger \in \mathcal{X}'$ which satisfies the Euler equation for G . Then $\lim_{t \rightarrow 0} x_t^\dagger = u^0$ and $\lim_{t \rightarrow \infty} x_t^\dagger = u^*$ by Lemma 4. It therefore suffices to show that x is a version of $x^\dagger$. We begin with a claim.

Claim. $\frac{1}{2}x + \frac{1}{2}x^\dagger$ belongs to $\mathcal{X}$, and $\pi_G\left(\frac{1}{2}x + \frac{1}{2}x^\dagger\right) \leq \frac{1}{2}\pi_G(x) + \frac{1}{2}\pi_G(x^\dagger)$.

Proof. x belongs to $\mathcal{X}$ by Lemma 0 (p. 14), since (x, X) is optimal (and optimality entails undominatedness by definition). Furthermore, $x^\dagger$ belongs to $\mathcal{X}' \subseteq \mathcal{X}$ by hypothesis. Since $\mathcal{X}$ is convex, it follows that $\frac{1}{2}x + \frac{1}{2}x^\dagger$ belongs to $\mathcal{X}$. For the remainder, x belongs to $\arg \max_{\mathcal{X}} \pi_G$ since (x, X) is optimal, and $x^\dagger$ belongs to $\arg \max_{\mathcal{X}} \pi_G$ by the Euler lemma. Hence

$$\frac{1}{2}\pi_G(x) + \frac{1}{2}\pi_G(x^\dagger) = \max_{\mathcal{X}} \pi_G \geq \pi_G\left(\frac{1}{2}x + \frac{1}{2}x^\dagger\right),$$

where the inequality holds since $\frac{1}{2}x + \frac{1}{2}x^\dagger$ belongs to $\mathcal{X}$. $\square$

Suppose toward a contradiction that x is not a version of $x^\dagger$. Since F^0, F^1, G are well-behaved, there are two cases to consider.

Case 1. F^0 is strictly concave on $[0, u^0]$. Choose a bounded non-null $A \subseteq \mathbf{R}_+$ on which $x \neq x^\dagger$, and note that $G(\sup A) < 1$ since G has unbounded support. Since F^0 is strictly concave on $[0, u^0]$ and F^1 is concave, it follows that $\pi_G\left(\frac{1}{2}x + \frac{1}{2}x^\dagger\right) > \frac{1}{2}\pi_G(x) + \frac{1}{2}\pi_G(x^\dagger)$, which contradicts the claim.

Case 2. F^1 is strictly concave on $[0, u^0]$ and G has full support. Since x is not a version of $x^\dagger$, there exists a bounded proper interval $I \subseteq \mathbf{R}_+$ on which $X \neq X^\dagger$. Since G has full support, I is G -non-null. Since F^0 (F^1) is concave (strictly concave) on $[0, u^0]$, it follows that $\pi_G\left(\frac{1}{2}x + \frac{1}{2}x^\dagger\right) > \frac{1}{2}\pi_G(x) + \frac{1}{2}\pi_G(x^\dagger)$, a contradiction with the claim. $\square$

F Proof of Proposition 3 (p. 24)

Recall the Euler equation defined in appendix D (p. 35).

Proposition 3'. Assume that F^0 and F^1 are differentiable on $(0, u^0)$ with bounded derivatives. Any mechanism that is optimal for G satisfies the Euler equation for G . Moreover, any undominated mechanism that satisfies the Euler equation for G is optimal for G .

This result refines Proposition 3 in two ways: it provides that the Euler equation is necessary under fewer assumptions, and furthermore asserts sufficiency.

Proof of Proposition 3'. Fix a distribution G . By Lemma 0 and Proposition 0 (pp. 14 and 15), any undominated mechanism has the form (x, X) with $x \in \mathcal{X}$. If (x, X) is undominated and satisfies the Euler equation for G , then it maximises the principal's payoff under G by the Euler lemma (appendix D, p. 35), so is optimal for G . Conversely, if (x, X) is optimal for G , then by the Euler lemma, (x, X) satisfies the Euler equation. $\square$

Proof of Proposition 3. Let (x, X) be optimal for a distribution G with $G(0) = 0$ and unbounded support. Then x is decreasing with $x \geq u^* > 0$ by Theorem 2 (p. 22), and (x, X) satisfies the Euler equation by Proposition 3'. Since G has unbounded support, Lemma 2 (appendix D, p. 36) yields that $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$ and that equation (5) holds for a.e. $t \in \mathbf{R}_+$. Since $F^{0'}$ is continuous and x decreasing, the right-continuous version of x satisfies (5) for all $t \in \mathbf{R}_+$. $\square$

Proposition 3' implies the assertion made in footnote 30 on p. 22:

Corollary 1. Let G be a distribution with unbounded support and $G(0) = 0$. Assume that F^0, F^1, G are well-behaved. Then any mechanism (x, X) that is optimal for G has $X_0 > u^1$.

Proof. If $u^* = u^1$, then $X_0 > u^* = u^1$ by Theorem 2 (p. 22). Assume for the remainder that $u^* < u^1$, and suppose toward a contradiction that $X_0 \leq u^1$. Then $X_t < u^1$ for all $t > 0$ since X is decreasing with $\lim_{t \rightarrow \infty} X_t = u^* < u^1$ by Theorem 2 (p. 22), and thus $X < u^1$ G -a.e. since $G(0) = 0$. Since F^1 is strictly increasing on $[0, u^1]$, it follows that $F^{1'}(X_s) > 0$ for G -a.e. $s \in \mathbf{R}_+$. Then since (x, X) satisfies the Euler equation by Proposition 3', it holds for a.e. $t \in \mathbf{R}_+$ with $G(t) > 0$ and $x_t < u^0$ that

$$0 < \int_{[0,t]} F^{1'}(X_s) G(ds) = -[1 - G(t)] F^{0'}(x_t) \leq 0,$$

which is absurd. $\square$

G Proof of Proposition 2 (p. 21)

As per appendix D, F^0 and F^1 are differentiable on $(0, u^0)$ with bounded derivatives, and we extend these continuously to $[0, u^0]$. (This means, in particular, that ' $F^{0'}(u^0)$ ' denotes the left-hand derivative of F^0 at u^0 .) Fix a distribution G . By inspection, for any deadline mechanism (x, X) , x belongs to $\mathcal{X}$ and satisfies (5) (appendix D, p. 36) for all $t \in \mathbf{R}_+$ with $G(t) < 1$, and furthermore $x \geq u^* > 0$. Hence by Lemma 2 (appendix D, p. 36), a deadline mechanism (x, X) satisfies the Euler equation for G iff $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$.

To prove the ‘only if’ part of Proposition 2, fix a mechanism (x, X) that is optimal for G . It is a deadline mechanism by Theorem 1 (p. 18), and satisfies the Euler equation for G by Proposition 3’ (appendix F, p. 43). Hence $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$ by the above.

To prove the ‘if’ part of Proposition 2, fix a deadline mechanism (x, X) such that $\mathbf{E}_G(F^{1'}(X_\tau)) = 0$. Since X is decreasing, it follows that $X_0 \geq u^1$, so (x, X) is undominated by Proposition 1 (p. 20). By the above, (x, X) satisfies the Euler equation for G . Hence by Proposition 3’, (x, X) is optimal for G . $\square$

References

- Acharya, V. V., DeMarzo, P., & Kremer, I. (2011). Endogenous information flows and the clustering of announcements. *American Economic Review*, 101(7), 2955–2979. <https://doi.org/10.1257/aer.101.7.2955>
- Armstrong, M., & Vickers, J. (2010). A model of delegated project choice. *Econometrica*, 78(1), 213–244. <https://doi.org/10.3982/ECTA7965>
- Atkeson, A., & Lucas, R. E., Jr. (1992). On efficient distribution with private information. *Review of Economic Studies*, 59(3), 427–453. <https://doi.org/10.2307/2297858>
- Atkeson, A., & Lucas, R. E., Jr. (1995). Efficiency and equality in a simple model of efficient unemployment insurance. *Journal of Economic Theory*, 66(1), 64–88. <https://doi.org/10.1006/jeth.1995.1032>
- Baron, D. P., & Besanko, D. (1984). Regulation and information in a continuing relationship. *Information Economics and Policy*, 1(3), 267–302. [https://doi.org/10.1016/0167-6245\(84\)90006-4](https://doi.org/10.1016/0167-6245(84)90006-4)
- Battaglini, M. (2005). Long-term contracting with Markovian consumers. *American Economic Review*, 95(3), 637–658. <https://doi.org/10.1257/0002828054201369>
- Ben-Porath, E., Dekel, E., & Lipman, B. L. (2019). Mechanisms with evidence: Commitment and robustness. *Econometrica*, 87(2), 529–566. <https://doi.org/10.3982/ECTA14991>
- Berkovitch, E., & Israel, R. (2004). Why the NPV criterion does not maximize NPV. *Review of Financial Studies*, 17(1), 239–255. <https://doi.org/10.1093/rfs/hhg023>
- Bird, D., & Frug, A. (2019). Dynamic non-monetary incentives. *American Economic Journal: Microeconomics*, 11(4), 111–150. <https://doi.org/10.1257/mic.20170025>
- Board, S. (2007). Selling options. *Journal of Economic Theory*, 136(1), 324–340. <https://doi.org/10.1016/j.jet.2006.08.005>

- Board, S., & Skrzypacz, A. (2016). Revenue management with forward-looking buyers. *Journal of Political Economy*, 124(4), 168–198. <https://doi.org/10.1086/686713>
- Boone, J., & van Ours, J. C. (2012). Why is there a spike in the job finding rate at benefit exhaustion? *De Economist*, 160(4), 413–438. <https://doi.org/10.1007/s10645-012-9187-8>
- Bull, J., & Watson, J. (2007). Hard evidence and mechanism design. *Games and Economic Behavior*, 58(1), 75–93. <https://doi.org/10.1016/j.geb.2006.03.003>
- Campbell, A., Ederer, F., & Spinnewijn, J. (2014). Delay and deadlines: Freeriding and information revelation in partnerships. *American Economic Journal: Microeconomics*, 6(2), 163–204. <https://doi.org/10.1257/mic.6.2.163>
- Courty, P., & Li, H. (2000). Sequential screening. *Review of Economic Studies*, 67(4), 697–717. <https://doi.org/10.1111/1467-937X.00150>
- Curello, G. (2023a). *Concealment in collective innovation* [working paper, 2 May 2023].
- Curello, G. (2023b). *Incentives for collective innovation* [working paper, 4 May 2023]. <https://doi.org/10.48550/arXiv.2109.01885>
- Curello, G., & Sinander, L. (2025). *Screening for breakthroughs* [working paper, 13 Mar 2025]. <https://doi.org/10.48550/arXiv.2011.10090>
- Czigler, T., Reiter, S., Schulze, P., & Somers, K. (2020). Laying the foundation for zero-carbon cement. *McKinsey & Company Chemicals Insights*. <https://www.mckinsey.com/industries/chemicals/our-insights/laying-the-foundation-for-zero-carbon-cement>
- de Clippel, G., Eliaz, K., Fershtman, D., & Rozen, K. (2021). On selecting the right agent. *Theoretical Economics*, 16(2), 381–402. <https://doi.org/10.3982/TE4027>
- DellaVigna, S., Lindner, A., Reizer, B., & Schmieder, J. F. (2017). Reference-dependent job search: Evidence from Hungary. *Quarterly Journal of Economics*, 132(4), 1969–2018. <https://doi.org/10.1093/qje/qjx015>
- Dilmé, F., & Li, F. (2019). Revenue management without commitment: Dynamic pricing and periodic flash sales. *Review of Economic Studies*, 86(5), 1999–2034. <https://doi.org/10.1093/restud/rdy073>
- Dye, R. A., & Sridhar, S. S. (1995). Industry-wide disclosure dynamics. *Journal of Accounting Research*, 33(1), 157–174. <https://doi.org/10.2307/2491297>
- Eső, P., & Szentes, B. (2007a). Optimal information disclosure in auctions and the handicap auction. *Review of Economic Studies*, 74(3), 705–731. <https://doi.org/10.1111/j.1467-937x.2007.00442.x>

- Eső, P., & Szentes, B. (2007b). The price of advice. *RAND Journal of Economics*, 38(4), 863–880. <https://doi.org/10.1111/j.0741-6261.2007.00116.x>
- Feng, F. Z., Taylor, C. R., Westerfield, M. M., & Zhang, F. (2024). Setbacks, shutdowns, and overruns. *Econometrica*, 92(3), 815–847. <https://doi.org/10.3982/ECTA21548>
- Frankel, A. (2016). Discounted quotas. *Journal of Economic Theory*, 166, 396–444. <https://doi.org/10.1016/j.jet.2016.08.001>
- Garrett, D. F. (2016). Intertemporal price discrimination: Dynamic arrivals and changing values. *American Economic Review*, 106(11), 3275–3299. <https://doi.org/10.1257/aer.20130564>
- Garrett, D. F. (2017). Dynamic mechanism design: Dynamic arrivals and changing values. *Games and Economic Behavior*, 104, 595–612. <https://doi.org/10.1016/j.geb.2017.04.005>
- Gershkov, A., Moldovanu, B., & Strack, P. (2018). Revenue-maximizing mechanisms with strategic customers and unknown, Markovian demand. *Management Science*, 64(5), 1975–2471. <https://doi.org/10.1287/mnsc.2017.2724>
- Glazer, J., & Rubinstein, A. (2004). On optimal rules of persuasion. *Econometrica*, 72(6), 1715–1736. <https://doi.org/10.1111/j.1468-0262.2004.00551.x>
- Glazer, J., & Rubinstein, A. (2006). A study in the pragmatics of persuasion: A game theoretical approach. *Theoretical Economics*, 1(4), 395–410.
- Green, B., & Taylor, C. R. (2016). Breakthroughs, deadlines, and self-reported progress: Contracting for multistage projects. *American Economic Review*, 106(12), 3660–3699. <https://doi.org/10.1257/aer.20151181>
- Grossman, S. J. (1981). The informational role of warranties and private disclosure about product quality. *Journal of Law & Economics*, 24(3), 461–483. <https://doi.org/10.1086/466995>
- Grossman, S. J., & Hart, O. D. (1980). Disclosure laws and takeover bids. *Journal of Finance*, 35(2), 323–334. <https://doi.org/10.2307/2327390>
- Guo, Y. (2016). Dynamic delegation of experimentation. *American Economic Review*, 106(8), 1969–2008. <https://doi.org/10.1257/aer.20141215>
- Guo, Y., & Hörner, J. (2020). *Dynamic allocation without money* [working paper, 26 Aug 2020].
- Guo, Y., & Shmaya, E. (2023). Regret-minimizing project choice. *Econometrica*, 91(5), 1567–1593. <https://doi.org/10.3982/ECTA20157>
- Guttman, I., Kremer, I., & Skrzypacz, A. (2014). Not only what but also when: A theory of dynamic voluntary disclosure. *American Economic Review*, 104(8), 2400–2420. <https://doi.org/10.1257/aer.104.8.2400>

- Hägele, I. (2024). *Talent hoarding in organizations* [working paper, Feb 2024]. <https://doi.org/10.48550/arXiv.2206.15098>
- Hansen, G. D., & İmrohoroğlu, A. (1992). The role of unemployment insurance in an economy with liquidity constraints and moral hazard. *Journal of Political Economy*, 100(1), 118–142. <https://doi.org/10.1086/261809>
- Hart, S., Kremer, I., & Perry, M. (2017). Evidence games: Truth and commitment. *American Economic Review*, 107(3), 690–713. <https://doi.org/10.1257/aer.20150913>
- Hopenhayn, H. A., & Nicolini, J. P. (1997). Optimal unemployment insurance. *Journal of Political Economy*, 105(2), 412–438. <https://doi.org/10.1086/262078>
- Jackson, M. O., & Sonnenschein, H. F. (2007). Overcoming incentive constraints by linking decisions. *Econometrica*, 75(1), 241–257. <https://doi.org/10.1111/j.1468-0262.2007.00737.x>
- Kyyrä, T., Pesola, H., & Verho, J. (2019). The spike at benefit exhaustion: The role of measurement error in benefit eligibility. *Labour Economics*, 60, 75–83. <https://doi.org/10.1016/j.labeco.2019.06.001>
- Li, J., Matouschek, N., & Powell, M. (2017). Power dynamics in organizations. *American Economic Journal: Microeconomics*, 9(1), 217–241. <https://doi.org/10.1257/mic.20150138>
- Lipnowski, E., & Ramos, J. (2020). Repeated delegation. *Journal of Economic Theory*, 188. <https://doi.org/10.1016/j.jet.2020.105040>
- Lyons, B. R. (2003). *Could politicians be more right than economists? A theory of merger standards* [working paper, Nov 2003].
- Madsen, E. (2022). Designing deadlines. *American Economic Review*, 112(3), 963–997. <https://doi.org/10.1257/aer.20200212>
- Matsushima, H., Miyazaki, K., & Yagi, N. (2010). Role of linking mechanisms in multitask agency with hidden information. *Journal of Economic Theory*, 145(6), 2241–2259. <https://doi.org/10.1016/j.jet.2010.03.014>
- Mierendorff, K. (2016). Optimal dynamic mechanism design with deadlines. *Journal of Economic Theory*, 161, 190–222. <https://doi.org/10.1016/j.jet.2015.10.007>
- Milgrom, P. (1981). Good news and bad news: Representation theorems and applications. *Bell Journal of Economics*, 12(2), 380–391. <https://doi.org/10.2307/3003562>
- Mirrlees, J. A. (1971). An exploration in the theory of optimum income taxation. *Review of Economic Studies*, 38(2), 175–208. <https://doi.org/10.2307/2296779>
- Mirrlees, J. A. (1974). Notes on welfare economics, information and uncertainty [reprinted in Mirrlees (2006)]. In M. S. Balch, D. McFadden,

- & S. Y. Wu (Eds.), *Essays on economic behavior under uncertainty* (pp. 243–258). North-Holland.
- Mirrlees, J. A. (2006). *Welfare, incentives, and taxation*. Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780198295211.001.0001>
- Nocke, V., & Whinston, M. D. (2013). Merger policy with merger choice. *American Economic Review*, 103(2), 1006–1033. <https://doi.org/10.1257/aer.103.2.1006>
- Pai, M. M., & Vohra, R. (2013). Optimal dynamic auctions and simple index rules. *Mathematics of Operations Research*, 38(4), 682–697. <https://doi.org/10.1287/moor.2013.0595>
- Pavan, A., Segal, I., & Toikka, J. (2014). Dynamic mechanism design: A Myersonian approach. *Econometrica*, 82(2), 601–653. <https://doi.org/10.3982/ECTA10269>
- Roberts, K. (1982). *Long term contracts* [working paper, Jul 1982].
- Rockafellar, R. T. (1970). *Convex analysis*. Princeton University Press.
- Rudin, W. (1976). *Principles of mathematical analysis* (3rd). McGraw-Hill.
- Sannikov, Y. (2008). A continuous-time version of the principal–agent problem. *Review of Economic Studies*, 75(3), 957–984. <https://doi.org/10.1111/j.1467-937X.2008.00486.x>
- Shavell, S., & Weiss, L. (1979). The optimal payment of unemployment insurance benefits over time. *Journal of Political Economy*, 87(6), 1347–1362.
- Sher, I. (2011). Credibility and determinism in a game of persuasion. *Games and Economic Behavior*, 71(2), 409–419. <https://doi.org/10.1016/j.geb.2010.05.008>
- Shimer, R., & Werning, I. (2008). Liquidity and insurance for the unemployed. *American Economic Review*, 98(5), 1922–1942. <https://doi.org/10.1257/aer.98.5.1922>
- Stein, E. M., & Shakarchi, R. (2005). *Real analysis: Measure theory, integration, and Hilbert spaces*. Princeton University Press.
- Thomas, J., & Worrall, T. (1990). Income fluctuation and asymmetric information: An example of a repeated principal–agent problem. *Journal of Economic Theory*, 51(2), 367–390. [https://doi.org/10.1016/0022-0531\(90\)90023-D](https://doi.org/10.1016/0022-0531(90)90023-D)
- Viscusi, W. K. (1978). A note on “lemons” markets with quality certification. *Bell Journal of Economics*, 9(1), 277–279. <https://doi.org/10.2307/3003627>