

Midterm Review

Doruk Cetemen* Yonggyun Kim[†] Fei Li[‡] Curtis R. Taylor[§]

September 7, 2026

Abstract

We study why organizations conduct interim performance reviews when monetary rewards are limited. An interim review creates incentive capacity by allowing future work and career opportunities to serve as rewards for past performance. Optimal review policies map a continuum of performance outcomes into a simple incentive ladder: termination, tough or easy continuation, and, for exceptional performance, an early maximal reward with no further work. Review can even sustain high effort when terminal compensation alone cannot. Its timing balances two forces: waiting improves the information revealed by performance, but leaves less future work available to motivate the agent.

JEL Classification: D82, D86, D23, M12.

Keywords: feedback, bounded reward, incentive ladder, incentive capital.

*Department of Economics and Financial Markets, Luiss Guido Carli University and EIEF: d cetemen@luiss.it

[†]Department of Economics, Florida State University: ykim22@fsu.edu

[‡]Department of Economics, University of North Carolina: lifei@email.unc.edu

[§]Department of Economics, Duke University: curtis.taylor@duke.edu

1 Introduction

Performance reviews are ubiquitous in organizations and often have substantial consequences for an employee’s career. Evaluations affect bonuses and merit pay, but also promotion, retention, and dismissal (Cappelli and Conyon, 2018). This raises a basic question. If an organization can ultimately observe an informative measure of performance and reward the employee accordingly, why evaluate performance before the end of the current phase of employment? In a standard moral-hazard model with unrestricted monetary transfers, there need be no reason to do so. Indeed, it is optimal to bundle incentives from different stages into a sufficiently powerful terminal compensation scheme.

This paper shows that the answer changes fundamentally when monetary rewards are bounded. Let M denote the largest performance-contingent payment that the principal can make. The cap may be literal, but more broadly it captures organizations in which large financial bonuses are difficult, costly, or institutionally constrained. Such restrictions are particularly natural in public-sector employment, where pay is relatively compressed and promotion is an important source of incentives (Mas, 2017; Deserranno et al., 2025), and in nonprofit organizations, which make less use of output-contingent incentive pay than for-profit firms (DeVaro and Brookshire, 2007). More generally, many careers provide substantial incentives through advancement rather than immediate cash compensation. Tenure and promotion in academia and promotion to partnership in law, accounting, consulting, and other professional-service firms are prominent examples (Baker et al., 1988; Barlevy and Neal, 2019; Levin and Tadelis, 2005). In these environments, the opportunity to continue on a favorable career path is itself a reward.

This observation gives interim review a role that is absent when monetary incentives are unbounded. A review allows the principal to allocate future opportunities according to past performance. An agent who performs well can be offered a more attractive continuation contract, while one who performs poorly can face a demanding continuation standard or termination. Rents that must be conceded to motivate future effort can therefore do double duty: besides inducing future effort, they can be used as rewards for past effort. We refer to these future rents and continuation opportunities as *incentive capital*. When M is large, the principal can instead provide incentives directly with money and incentive capital has little value. When M is tight, continuation opportunities become valuable precisely because monetary rewards are scarce.

We study this mechanism in a canonical dynamic moral-hazard environment. A risk-neutral agent produces cumulative Gaussian output over a finite horizon and privately chooses costly effort.¹ A risk-neutral principal receives output net of compensation. Compensation satisfies both limited liability and is bounded above by M . The principal may commit to one public interim review, choose when it occurs, and condition both compensation and the continuation opportunity on the review score. Past performance contains information about past effort but does not change the technology, productivity, or difficulty of future work. Thus review has no sorting, learning, or technological role in the baseline: any value it creates comes purely from incentives.

¹For ease of exposition, we assume that effort in the baseline model is chosen only twice, before the performance review and after it. We relax this assumption and allow the agent to choose effort continuously in Section 6.1.

Our first main result is that a continuum of possible review scores is optimally mapped into at most four discrete actions. These actions form an *incentive ladder*, ordered by the rent they deliver to the agent. At the bottom is a *bad stop*: sufficiently poor performance leads to termination with no payment. The next rung is continuation under a difficult terminal standard, which induces future effort while leaving the agent relatively little rent. Better review performance earns continuation under an easier terminal standard and hence a larger rent. Finally, when first-phase incentives are sufficiently scarce, exceptional performance reaches a *good stop*: the agent receives the maximal payment M immediately and is released from the remaining workload. Thus termination can optimally occur at both tails of the performance distribution. The bad stop is the strongest punishment available; the good stop is the strongest reward available once money alone can no longer reward performance. Depending on parameters, two, three, or all four rungs are used.

The discreteness of this ladder is notable because both performance and the feasible set of continuation contracts are continuous. Its structure comes from the geometry of bounded incentive provision. Among contracts that induce future effort, the principal uses only the minimum- and maximum-rent continuation contracts. Together with the two extreme no-effort actions—termination without pay and termination with the maximal payment—these are the only exposed actions in the principal’s statewise assignment problem. Monotone likelihood ratios then order the four actions by review performance. In this sense, a continuous performance evaluation optimally produces a coarse rating system with only a small number of economically distinct consequences.

Our second set of results establishes when review is valuable. The payment bound M is again the key comparative static. When M is sufficiently large, no genuine review can improve on the terminal contract: uniformly over all review dates, the principal strictly prefers to wait until the end. As the cap tightens, however, the terminal contract exhausts its ability to provide incentives and continuation rents become increasingly valuable. Near the smallest cap at which full effort can be implemented without review, an interior review is always strictly profitable and all four rungs of the incentive ladder are used.

Indeed, review can do more than reduce the cost of incentives. It can make high effort implementable when terminal evaluation alone cannot. At the no-review implementability boundary, splitting the project into two shorter incentive problems creates strict incentive slack in each phase. For payment caps slightly below that boundary, full effort is impossible under any terminal-only contract but can be induced with an interim review. Review therefore expands the organization’s effective incentive capacity: future rents that would otherwise be an unavoidable cost of motivating later effort become an additional currency with which to motivate earlier effort.

Our third main result concerns *when* the review should occur. Delay has an obvious informational advantage. With more performance accumulated before the review, the distributions generated by working and shirking are easier to distinguish. But waiting also uses up the very future opportunities that make review valuable. A later review leaves less remaining work, a smaller range of continuation rents, and therefore less incentive capital with which to reward first-phase performance. Optimal review timing balances better information against more time left to motivate.

The endpoint results make this tradeoff particularly sharp. An arbitrarily early review is strictly harmful: in the Brownian model the right derivative of review value at date zero is minus infinity. A very short probationary period generates almost no useful information, while bounded rewards prevent the principal from cheaply magnifying that information into incentives. Every strictly optimal review is therefore bounded away from the initial date. At the other endpoint, when M is tight, moving the review slightly before the terminal date has first-order value because even a short continuation phase creates scarce incentive capital. Thus the model predicts neither “review as soon as possible” nor “wait until performance is most informative.” The optimum is governed by the tension between learning about past effort and preserving future opportunities with which to reward it.

The optimal mechanism that ultimately awards a *prize* of either 0 or M also suggests a life-cycle interpretation. Interim evaluation should be especially valuable when an agent’s incentives depend heavily on future advancement—for example, for junior academics approaching tenure or young professionals competing for partnership or promotion. As careers mature and direct monetary compensation becomes a larger part of the feasible reward package, the incentive value of allocating future career opportunities should diminish. More generally, the model predicts that reviews should be most consequential precisely where organizations have limited scope to use large performance-contingent monetary payments.

Finally, we examine three departures from the baseline. To express the key mechanism in the simplest setup, the baseline lets the agent choose effort only once in each phase. Our first extension allows him to choose effort at every instant; the four-rung ladder continues to apply. Next, replacing Gaussian output with a general statistical experiment shows that the ladder does not depend on normality. Under suitable likelihood-ratio conditions, the same four exposed actions emerge, while the timing results depend on how rapidly useful evidence accumulates over short horizons. Finally, we show that random review timing can be valuable because it relaxes the sharp tradeoff between collecting more pre-review information and preserving post-review incentive capital. Together these extensions indicate that the four-rung incentive ladder is more robust than the baseline analysis alone suggests.

Related Literature This paper contributes to the literature on interim performance evaluation and dynamic incentive provision by asking three questions: Why review? How should review information be used? And when to review?

First, we identify a new rationale for interim review that arises from limited incentive provision. As in Lizzeri et al. (2002), we consider a setting in which past performance does not change the primitive continuation agency problem. With unrestricted monetary rewards, the principal can bundle incentives across the two phases into a single terminal contract, making review unnecessary. However, when payments are bounded, as in the static setting of Jewitt et al. (2008), review allows the principal to condition otherwise unavoidable continuation rents and the remaining workload on earlier performance, thereby creating history dependence and expanding incentive capacity. In

contrast, much of the literature considers environments in which earlier actions or outcomes affect the difficulty or return of subsequent work. In such settings, an important role of interim review is to reveal the severity of the continuation agency problem and allow the continuation contract to respond accordingly (see, e.g., Lukas, 2010; Chen and Chiu, 2013; Ely and Szydlowski, 2020). A separate rationale is efficient sorting: Ray (2007) uses interim evaluation to identify and terminate projects with poor continuation prospects. Relatedly, Liu (2026) studies resource allocation when the agent privately knows a project’s fixed quality and always prefers a larger allocation. A public midterm signal correlated with project quality allows the principal to discipline the agent’s initial report, effectively “levering” information from the review into the first-period allocation. Our setting has neither private information about project quality to elicit nor heterogeneous continuation prospects to uncover. Instead, review provides information about the agent’s hidden productive effort, and its value comes from using continuation decisions to provide incentives when monetary rewards are limited.

Second, we show that interim performance can be used through a simple, cutoff-based contract. Depending on interim performance, the agent is terminated or retained under a continuation contract that pays only following sufficiently good final performance. Termination can therefore occur at both tails: a bad stop punishes poor performance, while a good stop rewards strong performance by combining the maximal monetary payment with relief from future effort. Whereas Ely et al. (2025) jointly design dynamic incentives and the disclosure of performance feedback, our cutoff structure concerns how revealed interim performance is translated into termination and continuation incentives.

Good-performance termination also arises in other dynamic moral hazard settings (see, e.g., Spear and Wang (2005) and Sannikov (2008)), but for a different reason. In those models, favorable performance raises the promised value of a risk-averse agent until maintaining incentives becomes too costly, so termination occurs when continuation value reaches a sufficiently high level. Our good stop is instead a direct response to the payment cap. Once the maximal monetary prize has been reached, reducing the agent’s remaining workload provides the additional reward that money cannot. This logic is closer to Zhu (2013), where periods of shirking can optimally serve as relaxation periods that reward good performance.

Third, we identify a new tradeoff governing the timing of review. Delaying review has a familiar benefit: more accumulated performance makes the review signal more informative. But delay has a novel cost in our setting: it shortens the remaining production horizon and thereby reduces the future output that can be used to create continuation rents. Review timing therefore determines not only the quality of information available to the principal, but also the amount of incentive capacity remaining after the review.

This question connects our paper to a recent literature on the design and timing of dynamic monitoring. Varas et al. (2020) show that optimal monitoring can take the form of deterministic periodic reviews when information acquisition is important, but random inspections when monitoring primarily serves to provide incentives. Similarly, Ball and Knoepfle (2026) study whether inspec-

tions should be predictable or random in a long-term moral-hazard relationship, with the answer depending on whether effort primarily promotes breakthroughs or prevents breakdowns. Wong (2026) goes further by allowing the principal to design the monitoring technology itself, obtaining a nonstationary sequence of pass/fail tests whose informativeness and consequences evolve over the relationship. Relatedly, Dai et al. (2026) jointly design monitoring and compensation, with the principal dynamically reallocating limited monitoring capacity between evidence confirming effort and evidence revealing shirking. These papers endogenize the pattern, composition, or technology of monitoring, whereas we take a single interim review as given and ask where it should be placed within a finite production horizon.

Closest to our timing question, Rodivilov (2022) studies when to monitor an agent engaged in experimentation and shows that monitoring has both a static effect on contemporaneous moral hazard and a dynamic effect on the rents that must be provided in earlier periods. Our mechanism is different. In our model, the benefit of delay comes from the increasing informativeness of accumulated performance, while its cost comes from exhausting the very continuation opportunities through which review provides incentives.

This distinction also separates our mechanism from Hoffmann et al. (2021) and Georgiadis and Szentes (2020), who study the timing of performance measurement following a one-time hidden action rather than an interim review during ongoing production. In Hoffmann et al. (2021), delaying payout improves the information on which compensation is based but entails a direct cost of deferral; in Georgiadis and Szentes (2020), continued monitoring improves precision but incurs a direct monitoring cost. In our model, by contrast, the cost of delay arises from the contracting technology itself: waiting for a more informative review leaves less future work with which to reward favorable performance and provide incentives.

Organization The analysis proceeds as follows. Section 3 derives the no-review benchmark and constructs the post-review continuation frontier. Section 4 holds the review date fixed, formulates pre-review incentives, and characterizes the ordered menu. Section 5 endogenizes the date, deriving the review value, its behavior near the boundary dates, and the tight-cap result. Section 6 studies pointwise effort, non-Gaussian monitoring, and random review. Section 7 concludes with two directions requiring a richer institutional setting.

2 Environment

A risk-neutral principal (she) contracts with a risk-neutral agent (he) over a finite horizon $[0, T]$. Output X_t evolves as

$$dX_t = e_t dt + dW_t, \quad X_0 = 0,$$

where W is a standard Brownian motion and $e_t \in \{0, 1\}$ is agent effort. We write $\Phi(\cdot)$ for the standard normal distribution and $\phi(\cdot)$ for its density.

The contract fixes at date 0 a single review date $\tau \in (0, T]$.² The review date is ultimately a choice variable, and we solve for it in two steps: Section 4 takes τ as *exogenous*, and Section 5 then *endogenizes* it.

Assumption 1 (Observability of the output). *Cumulative output is publicly observed at τ and at T . No one, including the agent, observes X_t at any other time.*

The review splits the horizon into two phases, $P_1 = (0, \tau]$ and $P_2 = (\tau, T]$, of lengths τ and $T - \tau$; we call X_τ the *review score*. The case $\tau = T$ is the *no-review benchmark*: the second phase is empty and X_T is the only signal.

Assumption 2 (Phase-level private effort). *Effort is constant within each phase. The agent privately chooses $a_1 \in \{0, 1\}$ at date 0 and, having observed X_τ , privately chooses $a_2 \in \{0, 1\}$ at date τ ; effort is $e_t = a_1$ on P_1 and $e_t = a_2$ on P_2 . A strategy is a pair (a_1, S) with $S : \mathbb{R} \rightarrow \{0, 1\}$ giving $a_2 = S(X_\tau)$.³*

The agent commits at the start of a phase to work throughout it or shirk throughout it. Consequently the two phase increments are independent given (a_1, a_2) , with

$$X_\tau \sim N(a_1\tau, \tau), \quad X_T - X_\tau \sim N(a_2(T - \tau), T - \tau).$$

Working through a phase of length h costs the agent ch , so total effort cost is $c[a_1\tau + a_2(T - \tau)]$ with $c \in (0, 1)$. The same flow cost applies for the entire horizon, and hence in both phases. After observing X_τ and X_T the principal makes a measurable payment $w(X_\tau, X_T)$ to the agent.

Assumption 3 (Limited liability and finite compensation). *The payment function satisfies $w : \mathbb{R}^2 \rightarrow [0, M]$ with $M \in (0, \infty)$.*

The lower bound of 0 is the agent's limited liability, and the upper bound M is a finite cap: the maximum payment the principal can make.⁴

The principal receives ex post payoff $X_T - w$, and the agent receives $w - c[a_1\tau + a_2(T - \tau)]$. At $t = 0$ the principal chooses a contract that consists of a review date τ and a payment function w to maximize her expected payoff. Under the no-review benchmark $\tau = T$ a contract is simply $w(X_T) \in [0, M]$. The principal has all the bargaining power and full commitment, and the agent's outside option is zero.

3 Contracting over a single phase

We begin the analysis with the single-phase contract. Fix a phase of length $h > 0$ at the end of which output is observed and payment is made, with no further review. Two cases of the model

²The limiting case $\tau = 0$ is used only as a boundary limit in Section 5.

³Assumption 2 is relaxed in Section 6.1, where the agent chooses $e_t \in \{0, 1\}$ at each instant, so that a phase of length h delivers any total effort $\eta = \int e_t dt \in [0, h]$ rather than only $\{0, h\}$.

⁴All results generalize to a setting with finite lower and upper bounds on payments $-\infty < \underline{w} < \bar{w} < +\infty$, where $M \equiv \bar{w} - \underline{w}$. We set $\underline{w} = 0$ as is standard in settings of limited liability.

have this form, $h = T$ (no-review) and $h = T - \tau$ (post-review); in the latter, translation invariance of the technology lets a continuation contract be written as a function of the second phase output $Y = X_T - X_\tau$. Solving the phase problem once therefore delivers both cases.

Under work $Y \sim N(h, h)$ and under shirking $Y \sim N(0, h)$, with densities $f_j^h(y) = h^{-1/2} \phi((y - jh)/\sqrt{h})$ for $j \in \{0, 1\}$; working costs the agent ch . A contract is a measurable $w : \mathbb{R} \rightarrow [0, M]$; it implements work if

$$\mathbb{E}_1[w] - \mathbb{E}_0[w] \geq ch, \quad (\text{IC}_h)$$

where $\mathbb{E}_j[w] = \int w f_j^h$. Let $\mathcal{F}_h$ denote the set of contracts satisfying (IC_h) . Throughout the section c and h are fixed, and we suppress the dependence of derived objects on (h, c, M) except where an argument varies.

3.1 Threshold contracts

First, we show that it is without loss to focus on *threshold contracts*: pay the cap M when output clears a bar b and nothing below it. The reason is that the likelihood ratio $\ell^h(y) = f_0^h(y)/f_1^h(y) = \exp(h/2 - y)$ is strictly decreasing, so a contract that shifts payment toward high output — rewarding evidence of work and withholding payment after evidence of shirking — provides incentives at the lowest expected cost.

Lemma 1. *Let $w : \mathbb{R} \rightarrow [0, M]$ be any contract and let $\tilde{w} = \mathbf{1}_{\{y \geq b\}} M$ be the threshold contract with the same expected cost under work, $\mathbb{E}_1[\tilde{w}] = \mathbb{E}_1[w]$. Then $\mathbb{E}_0[\tilde{w}] \leq \mathbb{E}_0[w]$, with equality only if $w = \tilde{w}$ almost everywhere. Hence $w \in \mathcal{F}_h$ implies $\tilde{w} \in \mathcal{F}_h$.*

For any contract that implements work, there is a threshold contract that also implements work and leaves the principal the same expected payoff. Thus, nothing is lost by restricting attention to threshold contracts. Moreover, because every threshold contract pays either 0 or M , it can equivalently be interpreted as awarding a fixed prize M (e.g., a promotion) whenever performance exceeds a sufficiently high threshold.

Now, index the bar b by the *normalized* value $z = (h - b)/\sqrt{h}$, so that a larger z is an easier bar, cleared with probability $\Phi(z)$ under work. For $w = \mathbf{1}_{\{y \geq b\}} M$,

$$\mathbb{E}_1[w] = M\Phi(z), \quad \mathbb{E}_0[w] = M\Phi(z - \sqrt{h}), \quad \mathbb{E}_1[w] - \mathbb{E}_0[w] = M\Delta(h, z), \quad (1)$$

where $\Delta(h, z) = \Phi(z) - \Phi(z - \sqrt{h})$ is the *incentive index*: the increase in the probability of clearing the bar when the agent works rather than shirks. By (1) a threshold contract implements work exactly when

$$\Delta(h, z) \geq \frac{ch}{M}. \quad (2)$$

Figure 1 illustrates the admissible bars. The index is single-peaked at $z = \sqrt{h}/2$, where it attains its maximum $\bar{\Delta}(h) = 2\Phi(\sqrt{h}/2) - 1$, symmetric about that peak, and vanishes as $z \rightarrow \pm\infty$. Work is therefore implementable if and only if $M \geq M_I(h, c) \equiv \frac{ch}{\bar{\Delta}(h)}$, since below that cap (2) fails at

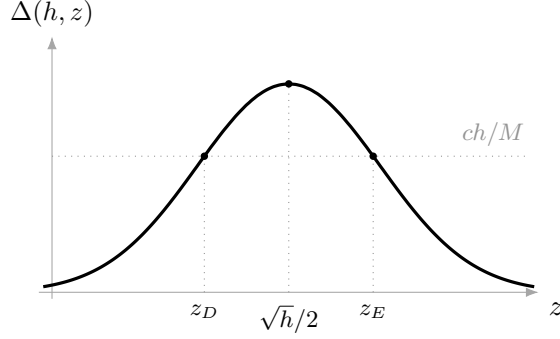


Figure 1: The incentive index $\Delta(h, \cdot)$

every z . When it holds, the admissible bars correspond to the closed interval $[z_D, z_E]$ with

$$z_D = \min \left\{ z : \Delta(h, z) \geq \frac{ch}{M} \right\}, \quad z_E = \max \left\{ z : \Delta(h, z) \geq \frac{ch}{M} \right\}, \quad (3)$$

the two crossings of the horizontal line at height ch/M , which merge at $\sqrt{h}/2$ when $M = M_I(h, c)$. The constraint binds at z_D and z_E and is slack strictly between them. The lower crossing z_D corresponds to a high-output bar and sets a *difficult* standard b_D ; the upper crossing z_E corresponds to a low bar and sets an *easy* one b_E .

Implementability does not guarantee that the principal wants the agent to work: the rent she must concede to do so may exceed what the phase produces. A second threshold therefore governs profitability.

Lemma 2. *There is a threshold $\underline{M}(h, c) \geq M_I(h, c)$ such that work is implementable and the principal's payoff from implementing it is nonnegative if and only if $M \geq \underline{M}(h, c)$. Additionally, there exists $\bar{c}(h) \in (0, 1)$ such that $\underline{M}(h, c) = M_I(h, c)$ for all $c \leq \bar{c}(h)$; otherwise $\underline{M}(h, c) > M_I(h, c)$.*

Appendix A gives $\bar{c}(h)$ and $\underline{M}(h, c)$ in closed form. We impose the following assumption to ensure that contracting without a review is strictly more profitable than shutting down.

Assumption 4 (Implementable and profitable full effort). *The payment cap satisfies $M > \underline{M}(T, c)$.*

3.2 The no-review benchmark

Setting $h = T$ gives the benchmark against which a review must be evaluated: the principal writes a single contract on terminal output and never intervenes. By Lemma 1 she may take it to be a threshold contract, and among the admissible bars she chooses the cheapest, which is the difficult one.

Proposition 1. *Under Assumption 4, the principal's optimal no-review payoff is*

$$V^{\text{NR}}(T; c, M) = T - M\Phi(z_D(T, c, M)) > 0,$$

attained by the difficult bar. Moreover $z_D(T, c, M) \rightarrow -\infty$ as $M \rightarrow \infty$, and

$$\lim_{M \rightarrow \infty} V^{\text{NR}}(T; c, M) = (1 - c)T,$$

so the first-best surplus accrues to the principal as the payment cap is relaxed.

The limit is what makes a review worth studying. Since $z_D = (T - b_D)/\sqrt{T}$, z_D falling to $-\infty$ corresponds to a bar b_D rising to $+\infty$ in output: as M grows the principal promises an ever larger prize on an ever less likely outcome. Because the Gaussian likelihood ratio is unbounded, that trade is favorable for the principal. The agent's rent is what he could secure by shirking, $\mathbb{E}_0[w] = M\Phi(z_D - \sqrt{T})$ by (1), and as the bar rises this vanishes even though M diverges: the reward event becomes overwhelming evidence of work. A principal who can promise arbitrarily large payments therefore captures the whole surplus without ever consulting interim performance, and has nothing to gain from a review.

When M is finite the agent retains a strictly positive rent, and the question is whether the principal can do better by making the continuation depend on an interim signal. The benchmark is generous to her in one further respect: under Assumptions 1 and 2 the agent commits at date 0 to work through the entire horizon and never observes his own progress, so he cannot condition effort on how the project is going.

3.3 The review action set

Now take $h = T - \tau$ to be a continuation phase. After a review score x the principal's choice is a continuation contract $w(x, \cdot) : \mathbb{R} \rightarrow [0, M]$ on the second-phase output Y , and the agent's best response to it determines whether second-phase work occurs. For everything that follows, such a contract is summarized by two numbers: the expected continuation utility r it leaves the agent net of effort cost, and the *total* expected continuation surplus s it generates—the principal's own expected continuation payoff is then $u = s - r$. We call the pair (r, s) an *action* and describe the review by the set of actions available at it. (We drop the qualifier *expected* from payoffs and surplus below whenever doing so creates no confusion.)

When effort is implemented, total surplus is $(1 - c)h$, so the only remaining question is which rents the agent can be promised.

Lemma 3. *The rents attainable while implementing work are exactly $[r_D, r_E]$, where*

$$r_j = M\Phi(z_j) - ch = M\Phi(z_j - \sqrt{h}), \quad j \in \{D, E\}, \quad 0 < r_D < r_E < M - ch,$$

and $r_D + r_E = M - ch$.

Here r_D and r_E are the rents left by the difficult and the easy standard respectively, and the lemma says that every attainable rent lies between them. The principal's payoff from the difficult standard, $V^{\text{NR}} = (1 - c)h - r_D$, is the no-review benchmark value when $h = T$. The set of rent-surplus pairs available under work is therefore $\mathcal{A}_1 = [r_D, r_E] \times \{(1 - c)h\}$.

If work is not implemented, surplus is zero, and since feasible payments range over $[0, M]$ the corresponding set is $\mathcal{A}_0 = [0, M] \times \{0\}$. The feasible action set is therefore $\mathcal{A} = \mathcal{A}_0 \cup \mathcal{A}_1$. A review menu assigns one such action to every review score. Thus, after each score, the principal chooses both a continuation profile—no work from $\mathcal{A}_0$ or work from $\mathcal{A}_1$ —and the rent delivered within that profile. The next section determines how these choices should vary with the review score in order to implement first-phase effort.

4 The exogenous review date

Having established what can be offered after the review, we now move backward to the incentives that precede it. Throughout this section the review date $\tau \in (0, T)$ is fixed exogenously, and we continue to write $h = T - \tau$.

4.1 The optimal incentive ladder

The review performance is $X_\tau \sim N(a\tau, \tau)$ if the agent chose action $e_t = a \in \{0, 1\}$ throughout the first phase. Let g_1^τ and g_0^τ be the corresponding densities. Their likelihood ratio $g_1^\tau(x)/g_0^\tau(x) = \exp(x - \tau/2)$ is strictly increasing.

Let $a(x) = (r(x), s(x)) \in \mathcal{A}$ be the action assigned after review score x , delivering the agent rent $r(x)$ and the principal payoff $s(x) - r(x)$. Working through the first phase yields the agent $\int r g_1^\tau - c\tau$ and shirking yields $\int r g_0^\tau$, so pre-review obedience requires

$$\int_{\mathbb{R}} r(x) [g_1^\tau(x) - g_0^\tau(x)] dx \geq c\tau, \quad (\text{OB}_\tau)$$

which also guarantees participation, since $r \geq 0$ makes the utility from shirking nonnegative. Conditional on inducing first-phase work, expected output accumulated by the exogenous review is τ , so the principal's problem at $t = 0$ is

$$\begin{aligned} V^R(\tau) = \sup_{a(\cdot)} \quad & \tau + \int_{\mathbb{R}} [s(x) - r(x)] g_1^\tau(x) dx \\ \text{subject to} \quad & (\text{OB}_\tau), \quad a(\cdot) \text{ measurable}, \quad a(x) \in \mathcal{A} \quad \text{for a.e. } x. \end{aligned} \quad (\text{P}_\tau)$$

Now we characterize how the continuation profile varies with the midterm review performance.

Proposition 2. *Fix $\tau \in (0, T)$ and suppose first-phase work is implemented. An optimal menu takes the ordered form*

$$a^*(x) = \begin{cases} A_B, & x < k_0^*, \\ A_D, & k_0^* \leq x < k_1^*, \\ A_E, & k_1^* \leq x < k_2^*, \\ A_G, & x \geq k_2^*, \end{cases}$$

where $A_B = (0, 0)$, $A_D = (r_D, (1 - c)h)$, $A_E = (r_E, (1 - c)h)$, $A_G = (M, 0)$ are four endpoints of $\mathcal{A}$, and $k_0^* \leq k_1^* \leq k_2^*$. Moreover, whenever both k_j^* and k_{j+1}^* are finite, the corresponding inequality is strict: $k_j^* < k_{j+1}^*$ for $j = 0, 1$.

Proposition 2 is our first main result. It shows that the optimal contract partitions the midterm-performance space into at most four ordered regions, assigning continuation profiles with progressively higher rents as performance improves. After the lowest performance realizations, A_B supplies the strongest punishment: the relationship ends with no payment and no second-phase production. At intermediate performance realizations the principal preserves the continuation surplus $(1 - c)h$. She first uses the difficult standard A_D , which delivers the smallest rent consistent with continuation effort, and then the easy standard A_E , which raises the agent's rent from r_D to r_E without sacrificing production. Only after the highest performance realizations can the principal move to A_G , paying the maximal rent M but giving up the continuation surplus altogether. This progression linking better review outcomes with higher discrete rent levels is what we refer to as the *incentive ladder*.

The optimal ladder therefore uses continuation contracts as intermediate incentive instruments. Before resorting to the maximal monetary reward, it increases the agent's continuation rent while keeping second-phase effort in place. The weak ordering of the cutoffs allows some of the upper profiles (rungs of the ladder) to be absent.

4.2 Deriving the optimal ladder

To understand why the optimal menu has this form, return to Problem (P_τ) . This is an optimization over the entire menu $a(\cdot)$, and its only constraint linking choices across review scores is pre-review obedience. Despite its infinite-dimensional appearance, the problem admits an elementary geometric solution. The reduction also reveals the economic forces governing the choice among stopping and continuation arrangements.

Let $\lambda^* \geq 0$ be a Lagrange multiplier on (OB_τ) at the optimum and define the *net shadow value of continuation rent*

$$Q_\lambda(x) \equiv \lambda \left(1 - \frac{g_0^\tau(x)}{g_1^\tau(x)} \right) - 1 = \lambda \left[1 - \exp\left(\frac{\tau}{2} - x\right) \right] - 1. \quad (4)$$

Adjoining (OB_τ) and factoring out $g_1^\tau(x)$ gives the Lagrangian

$$\mathcal{L}(a; \lambda, \tau) = \tau - \lambda c\tau + \int_{\mathbb{R}} g_1^\tau(x) \{s(x) + Q_\lambda(x) r(x)\} dx. \quad (5)$$

For each λ , the Lagrangian is maximized over measurable menus satisfying $a(x) \in \mathcal{A}$ for almost every x ; only pre-review obedience has been dualized. The optimal menu maximizes (5) at $\lambda = \lambda^*$, together with primal feasibility and complementary slackness. This reformulation separates the two agency problems. Post-review incentive compatibility and payment feasibility are encoded in the continuation-action set $\mathcal{A}$, while the multiplier incorporates pre-review obedience into the objective. Because the work-to-shirking likelihood ratio, and hence $Q_\lambda(x)$, increases with observed

performance, stronger performance places greater shadow value on the agent's continuation rent and therefore tilts the principal's choice toward higher-rent continuation profiles, subject to the surplus tradeoffs encoded by $\mathcal{A}$.⁵

For a fixed $\lambda > 0$, the Lagrangian objective is additively separable across review-score realizations. Choosing a measurable menu therefore reduces to solving, for almost every review score x , the linear assignment problem $s + Q_\lambda(x)r$. The coefficient $Q_\lambda(x)$ captures the weight placed on the agent's post-review continuation rent following score x . By Lemma B.1(iii) in Appendix B, the optimal supporting multiplier is positive outside a knife-edge case, so this characterization applies at the optimum.

Although the original action choice at a given review score is from $\mathcal{A} = \mathcal{A}_0 \cup \mathcal{A}_1$, it is convenient to solve the linear program

$$a_\lambda(x) \in \arg \max_{(r,s) \in \text{co}(\mathcal{A})} \{s + Q_\lambda(x)r\}.$$

Because the objective is linear in (r, s) , convexification does not change the maximal value, and an optimizer can always be selected from the original set $\mathcal{A}$.

A linear objective over the compact polygon $\text{co}(\mathcal{A})$ can always be maximized at an extreme point. Its four extreme points are A_B , A_D , A_E and A_G , all of which are defined in Proposition 2 and belong to $\mathcal{A}$. Hence, for each x , a solution to the above maximization problem can always be chosen from the original feasible set; public randomization cannot improve the principal's payoff. At a breakpoint, an entire edge of the convex hull may also be optimal, but one of its feasible endpoints attains the same value.

Figure 2 depicts this review-score assignment problem. The two bold segments form the original feasible set $\mathcal{A}$, while the shaded trapezoid is $\text{co}(\mathcal{A})$. The dashed lines are level sets of $s + qr$. Maximizing the linear objective amounts to moving such a line in the direction $(q, 1)$ until it supports the convex hull.

To solve the linear program, fix the scalar $q = Q_\lambda(x)$. The four vertices deliver

$$v_B(q) = 0, \quad v_D(q) = (1 - c)h + qr_D, \quad v_E(q) = (1 - c)h + qr_E, \quad v_G(q) = qM.$$

Their slopes are ordered by $0 < r_D < r_E < M$. The upper envelope therefore moves toward higher-rent actions as q rises. The adjacent crossings are

$$q_0 = -\frac{(1 - c)h}{r_D} < q_1 = 0 < q_2 = \frac{(1 - c)h}{M - r_E}. \quad (6)$$

Each breakpoint is an edge of Figure 2. A level set of $s + qr$ has slope $-q$ in (r, s) space, so at $q = q_j$ the supporting line is parallel to the edge joining the two actions that tie there: $-q_0 = (1 - c)h/r_D$ is the slope of $A_B A_D$, $-q_1 = 0$ that of the flat top edge $A_D A_E$, along which rent rises at no cost in surplus, and $-q_2 = -(1 - c)h/(M - r_E)$ that of $A_E A_G$. Since the upper boundary of $\text{co}(\mathcal{A})$ is

⁵This is the familiar likelihood-ratio formulation of a static moral-hazard problem (Holmström, 1979), except that the reward is a continuation profile rather than a wage.

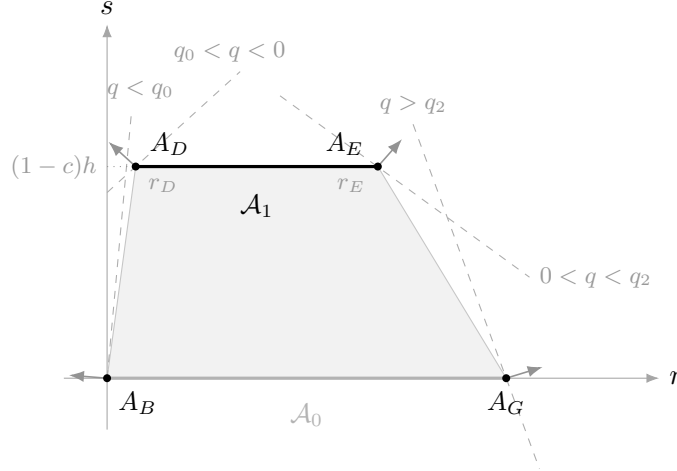


Figure 2: The review action set and the supporting lines of the assignment objective

concave, these three slopes fall from left to right, which is the ordering in (6). It follows that the maximizing action is A_B for $q < q_0$, A_D for $q_0 < q < q_1$, A_E for $q_1 < q < q_2$, and A_G for $q > q_2$, with adjacent actions tied at each breakpoint. Since $\lambda^* > 0$, $Q_{\lambda^*}(x)$ is strictly increasing in x . As performance rises, $q = Q_{\lambda^*}(x)$ therefore passes through these intervals in the same order. This converts the review-score ranking into the ordered performance regions stated in Proposition 2.

4.3 On the form of the optimal ladder

This subsection examines the economic forces that shape the specific form of the optimal menu—when termination is used following particularly poor or particularly strong performance, and when continuation is differentiated through more or less demanding standards.

Corollary 1. *Suppose $(1 - c)h \neq r_D$.⁶ Whenever first-phase work is implemented, k_0^* is finite and the bad-stop and difficult-standard regions are nonempty. Moreover:*

- (i) *If $1 + q_0 < \lambda^* \leq 1$, then $k_1^* = k_2^* = +\infty$, and only the bad-stop and difficult-standard regions are nonempty.*
- (ii) *If $1 < \lambda^* \leq 1 + q_2$, then $k_0^* < k_1^* < +\infty$ and $k_2^* = +\infty$, so the bad-stop, difficult-standard, and easy-standard regions are nonempty.*
- (iii) *If $\lambda^* > 1 + q_2$, then $k_0^* < k_1^* < k_2^* < +\infty$, and all four regions are nonempty.*

The multiplier λ^* is the shadow value of the first-phase obedience constraint. It measures the principal's marginal gain from relaxing the incentive requirement for first-phase work. A large multiplier therefore means that this requirement is particularly costly to satisfy: first-phase incentives are scarce, and the principal is willing to use stronger rewards following favorable performance.

⁶The knife-edge case $(1 - c)h = r_D$, in which a zero supporting multiplier may arise, is discussed in Appendix B.

This interpretation is reflected directly in the menu. The index $Q_{\lambda^*}(x)$ sweeps the interval $(-\infty, \lambda^* - 1)$ as performance x ranges over $\mathbb{R}$, and it crosses breakpoint q_j if and only if $\lambda^* > 1 + q_j$. Implementing first-phase work requires $\lambda^* > 1 + q_0$, so the bad stop and the difficult standard are always used. Once $\lambda^* > 1$, the incentive value of additional rent is high enough to justify the easy standard, which rewards stronger performance while preserving second-phase production. Once $\lambda^* > 1 + q_2$, even the easy standard provides too little rent at the top: the principal then uses the good stop, paying the maximal reward despite sacrificing the continuation surplus.

The range of Q_{λ^*} also explains the asymmetry between the lower and upper regions. It is unbounded below, so sufficiently poor performance always receives the strongest punishment. Its upper bound is $\lambda^* - 1$, however, so the stronger rewards become active only when the shadow value of first-phase incentives is sufficiently large. Whenever a breakpoint is reached, the corresponding cutoff is the unique solution to $Q_{\lambda^*}(k_j^*) = q_j$.

An especially sharp implication emerges when the review occurs late. As τ approaches T , the remaining continuation horizon becomes short. The continuation rungs do not disappear, however. Instead, the threshold for activating the highest reward converges to its lower bound, while the shadow value of first-phase incentives remains strictly above that bound. Consequently, every rung of the incentive ladder is used sufficiently close to the terminal date.

Proposition 3. *Suppose Assumption 4 holds. Then $k_2^*(\tau) < +\infty$ for every review date τ sufficiently close to T from below.*

For any finite admissible cap, a sufficiently late review uses all four continuation profiles. To see the force behind this result, recall that the good stop becomes active when $\lambda^* > 1 + q_2$. As the continuation horizon $h = T - \tau$ vanishes, the surplus sacrificed by replacing the easy standard with the good stop shrinks faster than the additional rent thereby created. Hence $q_2(h) \rightarrow 0$. At the same time, the fixed-review shadow value converges to the no-review shadow value, which is strictly greater than one under Assumption 4. Thus the good-stop threshold is eventually crossed, and all four regions are active.

5 The timing and value of review

In this section, the principal chooses the review date endogenously from $\tau \in (0, T]$. A date $\tau \in (0, T)$ represents a genuine interim review, whereas $\tau = T$ represents no interim review. A review at $\tau = 0$ is equivalent to one at $\tau = T$, so we use $\tau = T$ to represent no review and do not consider $\tau = 0$ separately. We first study how the value of a genuine review varies with its date and then determine when such a review dominates no review.

5.1 The optimal genuine review date

For the remainder of the timing analysis, fix a genuine review date $\tau \in (0, T)$ and let $h = T - \tau > 0$. Using the fixed-date characterization from Section 4, define the continuation support function for a

remaining horizon h and shadow value q by

$$\mathcal{V}(q, h) \equiv \max \{0, (1 - c)h + q r_D, (1 - c)h + q r_E, qM\}.$$

Complementary slackness implies that the optimized value equals the Lagrangian at the constrained optimum. Hence

$$V^R(\tau) = \tau - \lambda^* c \tau + \int_{\mathbb{R}} g_1^\tau(x) \mathcal{V}(Q_{\lambda^*}(x), T - \tau) dx,$$

where $Q_{\lambda^*}(x)$ is the net shadow value of continuation rent defined in Equation (4).

Lemma 4. *Suppose Assumption 4 holds. Then*

$$\lim_{\tau \downarrow 0} V^R(\tau) = \lim_{\tau \uparrow T} V^R(\tau) = V^{\text{NR}}(T; c, M).$$

Both boundaries degenerate, for opposite reasons. As $\tau \downarrow 0$ the review score carries almost no information about first-phase effort, so there is nothing to condition on; as $\tau \uparrow T$ the continuation phase vanishes, so there is no continuation opportunity left to allocate. Either way the principal is left with the terminal contract. The no-review benchmark is therefore the right normalization across the whole interval, and any gain from review must be an interior phenomenon.

The tradeoff among genuine review dates. Away from knife-edge parameters, the constrained envelope theorem applies while holding the multiplier and the selected action fixed. The derivative decomposes as

$$\begin{aligned} \frac{dV^R(\tau)}{d\tau} = & \underbrace{1 - \lambda^* c}_{\text{first-phase production}} + \underbrace{\int_{\mathbb{R}} \left[\frac{\partial g_1^\tau(x)}{\partial \tau} \mathcal{V}(Q_{\lambda^*}(x), h) + g_1^\tau(x) \frac{\partial \mathcal{V}}{\partial q}(Q_{\lambda^*}(x), h) \frac{\partial Q_{\lambda^*}(x)}{\partial \tau} \right] dx}_{\text{information}} \\ & - \underbrace{\int_{\mathbb{R}} g_1^\tau(x) \frac{\partial \mathcal{V}}{\partial h}(Q_{\lambda^*}(x), h) dx}_{\text{incentive capital}}. \end{aligned} \quad (7)$$

The first term captures first-phase production. Delaying the review lengthens the first phase, adding a unit of production while tightening first-phase obedience, whose shadow value is charged at rate c . The second term is the information effect. It values a sharper performance measure: the review-score means under working and shirking differ by τ against a standard deviation $\sqrt{\tau}$, so their standardized separation grows with the review date; since the marginal value of the index is exactly the rent the menu assigns, precision is worth most at scores that already carry a large rent. The third term is the opposing loss of incentive capital. A longer continuation phase is *incentive capital*: a family of work-inducing contracts that produce second-phase output while delivering different rents, so that rents conceded for second-phase incentives can be recycled into rewards for first-phase effort. As h falls this family shrinks, and strong first-phase rewards increasingly require sacrificing second-phase production, culminating in A_G .

An interior optimum balances these three terms, setting (7) equal to zero. The optimal review date therefore does not merely trade an early noisy signal against a late precise one; it balances the precision of the review against the productive contracting capacity that survives it.

The suboptimality of extreme review dates. The next result distinguishes the implications of the review-date tradeoff near the two boundaries.

Proposition 4. *Review dates sufficiently close to the initial date are strictly suboptimal. Review dates sufficiently close to the terminal date are also strictly suboptimal when the payment cap is sufficiently tight. Formally,*

- (i) *for every τ sufficiently close to the initial date, there exists another genuine review date τ' such that $V^R(\tau') > V^R(\tau)$;*
- (ii) *suppose $c \leq \bar{c}(T)$, so that $\underline{M}(T, c) = M_I(T, c)$. If M is sufficiently close to $\underline{M}(T, c)$ from above, then every review date τ sufficiently close to the terminal date is strictly dominated by an earlier genuine review date $\tau' < \tau$.*

The first conclusion reflects the information term in (7). A short first phase adds only a small amount of production and generates an extremely weak signal of effort. To motivate first-phase effort, the principal must nevertheless create enough dispersion in continuation rents across review outcomes. Because those rents are bounded, extracting the required incentives from an almost uninformative signal is disproportionately costly. The first-phase production gain therefore cannot compensate for the incentive loss generated by a review sufficiently close to the initial date.

The second conclusion reflects the opposing incentive-capital term. Near the terminal date, delaying the review yields a more informative score but leaves almost no continuation phase over which rents can be varied. A tight payment cap makes this loss especially costly: the terminal contract is already close to exhausting its incentive capacity, so continuation rents have a high shadow value. The marginal loss of incentive capital then dominates the production and information benefits of further delay, and an earlier review raises the principal's payoff.

As a result, under the tight-cap conditions of Proposition 4, if an optimal genuine review exists, it cannot lie near either the beginning or the end of the project. It is genuinely a *midterm* review. An implication is that short probationary periods observed in practice must therefore reflect motives absent from the model, such as screening an unknown type, learning match quality, or preserving an option to abandon a bad project.

5.2 When is a genuine review valuable?

We now compare genuine review with the outside option of no review. Define the review gain for $\tau \in (0, T)$ by

$$G(\tau) = V^R(\tau) - V^{\text{NR}}(T; c, M).$$

Loose payment caps. When the payment cap is loose, the no-review contract can concentrate compensation on an increasingly rare terminal outcome and implement effort with vanishing rent. A review cannot replicate this economy without either splitting the long experiment or relying on an almost uninformative early signal.

Proposition 5. *Fix $T > 0$ and $c \in (0, 1)$. There exists a finite $\overline{M}(T, c)$ such that*

$$G(\tau) < 0 \quad \text{for every } \tau \in (0, T)$$

whenever $M > \overline{M}(T, c)$.

The result is uniform over review dates. A pointwise comparison would still allow the principal to move the review date as M rises and stay ahead of the benchmark. Proposition 5 rules this out. The rent compression established in Proposition 1 eventually dominates uniformly: no allocation of production between a review phase and a continuation phase can preserve both the precision and the production horizon of the terminal contract.

Tight payment caps. The comparison can reverse when the payment cap is sufficiently close to the implementability boundary. The next result records the resulting local value improvement.

Proposition 6. *Suppose $c \leq \bar{c}(T)$, so that $\underline{M}(T, c) = M_I(T, c)$. If M is sufficiently close to $\underline{M}(T, c)$ from above, there exists $\delta > 0$ such that*

$$V^R(\tau) > V^{\text{NR}}(T; c, M) \quad \text{for every } \tau \in (T - \delta, T).$$

The logic combines the boundary value with the local comparison in Proposition 4. By Lemma 4, $G(\tau)$ converges to zero as τ approaches T from below. Under the stated tight-cap conditions, Proposition 4 shows that moving a terminal-adjacent review earlier strictly raises $G(\tau)$. Together, these observations imply that $G(\tau) > 0$ for review dates sufficiently close to T , as stated in Proposition 6.

Economically, moving the review slightly earlier creates a short continuation phase. When the cap is tight, the no-review contract is close to exhausting its incentive capacity, so the rents needed to induce continuation effort have a high shadow value and can also reward first-phase performance. This is what makes a sufficiently late review valuable. Separately, Proposition 3 shows that, under the maintained assumptions, the optimal fixed-date contract uses the full incentive ladder whenever the review is sufficiently close to T , regardless of how tight the payment cap is. Hence, in the tight-cap region covered by Proposition 6, the value improvement is implemented through all four incentive margins characterized in Section 4.3. The resulting gain in incentive capacity outweighs the small losses in information and production.

Payment caps and the value of review. The preceding results show how the payment cap governs the value of review. A review expands incentive capacity, but relies on a shorter performance

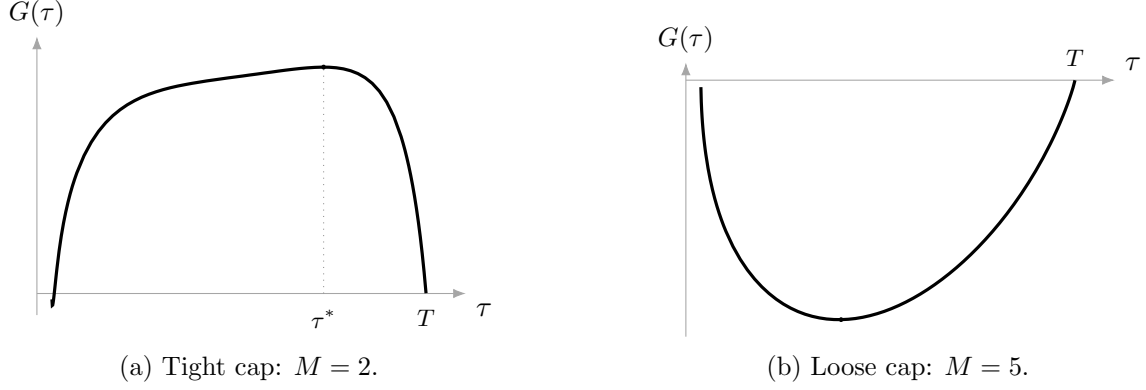


Figure 3: The review gain $G(\tau)$ under a tight and a loose payment cap

record and may forgo continuation output after a stop. A tight cap raises the value of the additional incentive margins; a loose cap lets the terminal contract provide incentives at negligible rent, leaving the information and production costs of review to dominate.

Figure 3 illustrates the two regimes. Both panels hold $T = 1.5$ and $c = 0.6$ fixed and vary only the payment cap, so the contrast is a comparative static in M alone. Under the tight cap $M = 2$, just above the implementability floor $M_I(T, c) = 1.96$, the gain is negative at the earliest review dates, turns positive, and reaches a strictly positive interior maximum at $\tau^* = 1.09$, where $V^R = 0.239$ against $V^{\text{NR}} = 0.195$. The supporting multiplier there is $\lambda^* = 2.19 > 1 + q_2$, so by Corollary 1 the optimal menu uses all four rungs of the ladder. Under the loose cap $M = 5$ the gain is negative at every review date, and the principal optimally waits until T . In both panels G vanishes at the two boundary dates, as Lemma 4 requires. The short initial interval on which $G < 0$ in the left panel is the finite- τ counterpart of the infinitely steep right derivative at zero in Proposition 4(i): an almost uninformative probation cannot pay for the rent dispersion it requires.

Review restores implementability under tight caps. Under tight caps, the value of a midterm review also lies in restoring implementability. At the no-review implementability boundary, splitting the horizon creates strict incentive slack in both phases, and a review permits full effort for caps slightly below the level required without one.

Proposition 7. *Let $M_0 \equiv M_I(T, c)$, at which full effort without review is just implementable. For every split $\tau \in (0, T)$, with $h = T - \tau$,*

$$M_0 \overline{\Delta}(\tau) > c\tau, \quad M_0 \overline{\Delta}(h) > ch.$$

Moreover there exist $\tau_0 \in (0, T)$ and $\varepsilon > 0$ such that for every $M \in (M_0 - \varepsilon, M_0)$ full effort is infeasible without a review, but a review at τ_0 admits a two-region menu inducing first-phase work and, whenever the project is continued, full continuation effort. Continuation occurs with positive probability.

A review can therefore create incentives unavailable when performance is evaluated only at

the terminal date. Splitting the project gives two shorter incentive problems, each strictly easier to implement from $M_0\bar{\Delta}(\tau) > c\tau$ and $M_0\bar{\Delta}(h) > ch$. More importantly, it creates a new source of motivation: the continuation phase has a nondegenerate range of work-inducing rents $[r_D, r_E]$, and the principal can condition those rents on the review score. Rents that must in any event be conceded to motivate future effort thus double as rewards for past effort, supplemented when necessary by the two stopping actions.

6 Extensions

The baseline model isolates the role of an endogenous interim review under phase-level effort, Gaussian progress, and deterministic review. This section summarizes three extensions that relax those assumptions in turn. The main text records the economically important objects and the structural conclusions; derivations, complete statements, proofs, and numerical illustrations are collected in the Online Appendix. A useful theme across the extensions is that the four-rung ladder is more robust than the exact Gaussian formulas: what changes first is typically the location or availability of a rung, while the ordering logic survives whenever the review score continues to rank continuation rents monotonically.

6.1 Pointwise effort choice

The baseline model restricts the agent to a phase-by-phase effort choice: in a phase of length h , he either works throughout or does not work at all. This section allows him to choose effort pointwise. Because no new information arrives within a phase, only total effort matters. Writing $e_t \in \{0, 1\}$ for instantaneous effort, let

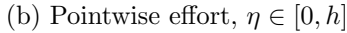
$$\eta = \int_0^h e_t dt \in [0, h].$$

Total effort η produces $Y \sim N(\eta, h)$ and costs $c\eta$. Thus the binary choice $\{0, h\}$ in the baseline becomes the interval $[0, h]$: the agent can now “coast” by working for only part of the phase.

This extension has two sharp implications. First, the four-rung structure of the optimal review menu survives: no interior effort level becomes an exposed continuation action. Second, and more importantly, the easy continuation standard must become harder. Pointwise effort leaves the difficult continuation unchanged, but reduces the largest rent that the principal can deliver through the easy continuation. The extension therefore preserves the form of the incentive ladder while weakening one of its rungs.

From a global to a marginal incentive test. For a terminal wage w , let $\pi(\eta) = \mathbb{E}_\eta[w(Y)]$ denote expected compensation when total effort is η . A contract implements the target η exactly when no other effort level is more attractive:

$$\pi(\eta) - c\eta \geq \pi(\eta') - c\eta' \quad \text{for every } \eta' \in [0, h]. \quad (\text{IC}_h^{\text{flex}})$$



The phase-by-phase model checks only the extreme deviation $\eta' = 0$. Pointwise choice also introduces nearby ones, including the temptation to shave the last increment of effort. Two consequences of $(\text{IC}_h^{\text{flex}})$ carry the analysis: the comparison with $\eta' = 0$, and local optimality,

$$\pi'(\eta) = c \text{ if } \eta \in (0, h), \quad \pi'(h) \geq c \text{ if } \eta = h. \quad (\text{L}_\eta)$$

The comparison with phase-by-phase effort is now immediate. At the difficult bar $z_D(h)$, the global condition binds and the marginal condition is slack. The difficult continuation is therefore unchanged. At the easy bar $z_E(h)$, by contrast, the baseline global condition binds but the marginal return to the last increment of effort is too small. The reason is simple: an easy bar is cleared with high probability even if the agent works slightly less, so coasting barely changes the probability of payment while saving effort cost. A contract can therefore induce “work rather than shirk” and still fail to induce “work all the way.”

$$\sqrt{h}\phi(\zeta(h)) = \frac{ch}{M}, \quad \text{equivalently} \quad \phi(\zeta(h)) = \frac{c\sqrt{h}}{M}, \quad (8)$$
$$\tilde{A}_E = (\tilde{r}_E(h), (1-c)h), \quad \tilde{r}_E(h) = M\Phi(\zeta(h)) - ch < r_E(h). \quad (9)$$

20

Figure 4 displays the central comparison. Moving from phase-by-phase to pointwise effort fills out the feasible rent–surplus set, but its convex hull still has four relevant vertices. Three are unchanged: bad stop A_B , difficult continuation A_D , and good stop A_G . Only the maximum-rent full-effort endpoint moves, from A_E to $\tilde{A}_E$. Geometrically, the full-effort face is shortened; economically, the easy continuation loses incentive capital.

The four-rung menu survives. Online Appendix OA.1 proves that the convex hull of $\mathcal{A}^{\text{flex}}(h)$ is $\text{co}\{A_B, A_D, \tilde{A}_E, A_G\}$: although pointwise effort permits a continuum of targets, no interior $\eta \in (0, h)$ maximizes the principal’s statewise linear objective. The optimal menu can therefore still be chosen from four actions, ordered by review performance exactly as in the baseline—bad stop, difficult continuation, easy continuation, good stop—and all baseline formulas carry over after the economically meaningful substitution $r_E(h) \mapsto \tilde{r}_E(h)$. In particular the upper breakpoint becomes $\tilde{q}_2(h) = (1 - c)h / (M - \tilde{r}_E(h))$, while the lower breakpoint and the ordering of the four regions are unchanged. Since $\tilde{r}_E(h) < r_E(h)$, the easy continuation is a weaker reward: the principal uses it less and relies relatively more on the indivisible good-stop prize. Pointwise effort therefore changes the strength, but not the architecture, of the incentive ladder. This substitution away from continuation rents toward good stopping is costly: it reduces productive continuation and lowers the principal’s payoff.

6.2 Beyond Gaussian noise

The normal distribution delivers closed-form thresholds, but it is not what creates the ordered incentive ladder. To separate the distribution-free argument from the Gaussian formulas, replace the Brownian increment over a phase of length h by a public signal with law P_1^h under work and law P_0^h under shirking. We retain stationary independent increments, so the continuation experiment depends on the review date only through the remaining horizon. Let $m(h)$ be the incremental expected output from work and define continuation surplus net of effort cost by $\Sigma(h) = m(h) - ch$.

From threshold contracts to statistical tests. In the Gaussian baseline, the monotone likelihood ratio makes the optimal wage a cutoff in output. With a general signal there need not be a natural output cutoff. Instead, normalize the wage as $\varphi = w/M \in [0, 1]$. A fractional wage $M\varphi(y)$ has the same expected payoffs as paying the cap with probability $\varphi(y)$ after signal y . The contract can therefore be represented as a possibly randomized statistical test: its *size* $\alpha = \mathbb{E}_0[\varphi]$ is the probability of payment under shirking, and its *power* $\beta = \mathbb{E}_1[\varphi]$ is the probability of payment under work.

Let $\mathcal{R}_h$ be the set of all feasible size–power pairs (α, β) , and let $\beta_h(\alpha)$ be the greatest power attainable at size α . The incentive spread from a pair is $M(\beta - \alpha)$. Its largest possible value is $M \text{TV}_h$, where $\text{TV}_h = \max_{(\alpha, \beta) \in \mathcal{R}_h} (\beta - \alpha)$, is the total-variation distance between the work and shirk signal laws, so the one-phase feasibility condition is exactly $ch \leq M \text{TV}_h$. Under strict feasibility, the work-inducing tests have sizes in the interval whose endpoints are $\alpha_L(h) = \min\{\alpha : \beta_h(\alpha) - \alpha \geq ch/M\}$

and $\alpha_H(h) = \max\{\alpha : \beta_h(\alpha) - \alpha \geq ch/M\}$. The corresponding work-inducing rents are exactly the interval $[r_L(h), r_H(h)] = M[\alpha_L(h), \alpha_H(h)]$, and the full-effort frontier is the horizontal segment $s = \Sigma(h)$ over that interval. As in the baseline, the principal's payoff $u = \Sigma(h) - r$ falls one for one with the rent conceded.

Two-sided evidence in both phases, slack detection, positive continuation surplus, monotone review evidence, and public randomization at score atoms recover the baseline ladder. Under these conditions $0 < r_L < r_H < M$, and the support function for the statewise objective $s + qr$ is

$$\mathcal{V}(q, h) = \max\{0, \Sigma(h) + qr_L(h), \Sigma(h) + qr_H(h), qM\}.$$

The four exposed actions are bad stop $A_B = (0, 0)$, low-rent continuation $A_L = (r_L, \Sigma)$, high-rent continuation $A_H = (r_H, \Sigma)$, and good stop $A_G = (M, 0)$. A monotone likelihood ratio for the review score assigns them in that order as performance improves. Thus Gaussianity is useful because it converts likelihood-ratio tests into signal cutoffs and supplies closed-form formulas; it is not needed for the four-action geometry.

Rent is cost divided by evidence. The general experiment also makes the source of agency rent transparent. If $\bar{\Lambda}_h$ is the average work-to-shirk likelihood ratio on the reward event of the cheapest work-inducing test, then $r_L(h) = ch/(\bar{\Lambda}_h - 1)$: rent is effort cost divided by the excess evidentiary content of the reward event. Consequently, a loose cap eliminates rent only if rare reward events can carry arbitrarily strong evidence of work. Under the regularity conditions in the Online Appendix, $\lim_{M \rightarrow \infty} r_L(h) = ch/(\Lambda_h^{\max} - 1)$, where $\Lambda_h^{\max}$ is the essential supremum of the work-to-shirk likelihood ratio, with the right-hand side interpreted as zero when $\Lambda_h^{\max} = \infty$. Bounded evidence can therefore leave a positive rent floor even as the payment cap becomes arbitrarily loose.

What can fail. The five conditions isolate distinct departures from the baseline. One-sided evidence can push an extreme rent to 0 or M and remove a rung. Nonpositive continuation surplus can make stopping technologically attractive rather than purely incentive driven. Failure of monotone likelihood ratios can destroy the ordering in the observed score, while atoms can require public mixing at a review threshold. The two likelihood-ratio tails also govern whether the stopping rungs are reached: bounded evidence of work makes the good stop harder to use, and bounded evidence of shirking similarly limits the bad stop.

Timing depends on short-horizon evidence. The review-date results require more than the feasible size–power set. Brownian evidence becomes locally weak as the phase shrinks, which produces the infinitely negative right derivative at $\tau = 0$.⁷ A rare event whose likelihood ratio remains bounded away from one over a short interval can instead provide incentives at first order.

⁷This distinction is related to Li (2017), who shows that review contracts can attain near-efficiency under Poisson and Gamma monitoring but not under Brownian monitoring. Under Brownian monitoring, short performance records remain too noisy to support sufficiently demanding performance standards.

Such hard evidence can remove the Brownian infinite-slope conclusion, although it does not by itself make an early review optimal. It can also give a vanishing continuation phase first-order value near T , breaking the Brownian equivalence between the review derivative and the no-review derivative. The distribution-free conclusion is thus a separation: the feasible size–power set delivers the ladder, whereas likelihood-ratio tails and short-horizon asymptotics determine which rungs are reached and how review timing behaves. Online Appendix OA.2 states the exact assumptions and derives the generalized endpoint formulas.

6.3 Random review dates

The review date has so far been announced at date 0. This extension lets the principal commit instead to a distribution Γ over review dates on $[0, T]$, draw the date independently of the Brownian motion, and withhold its realization until the review occurs. She commits publicly to Γ and to a date-contingent menu $a_\tau(\cdot)$ for every realization τ ; only the draw is hidden.

Timing is valuable for two conflicting reasons. A later review aggregates more performance and so yields a sharper score; an earlier one leaves a longer continuation phase and so more incentive capital. A deterministic date must settle on one mixture of the two. A lottery need not, because the agent chooses pre-review effort without knowing which date will arrive.

Phase-level effort. Under Assumption 2 the agent chooses a_1 once, at date 0, and cannot revise it before the review. Non-arrival therefore changes his beliefs but gives him no new action, and obedience is a single scalar constraint. Writing

$$P(\tau, a_\tau) = \tau + \int_{\mathbb{R}} [s_\tau(x) - r_\tau(x)] g_1^\tau(x) dx, \quad J(\tau, a_\tau) = \int_{\mathbb{R}} r_\tau(x) [g_1^\tau(x) - g_0^\tau(x)] dx - c\tau \quad (10)$$

for the principal’s payoff and the incentive surplus of a date–menu pair, the problem is

$$V^S = \sup_{\Gamma, \{a_\tau\}} \int P(\tau, a_\tau) \Gamma(d\tau) \quad \text{subject to} \quad \int J(\tau, a_\tau) \Gamma(d\tau) \geq 0. \quad (\text{SR})$$

Thus $J \geq 0$ is precisely the obedience constraint of the deterministic problem, and $V^D \equiv \sup_\tau V^R(\tau)$ is the value when Γ is degenerate. At $\tau = T$ there is no continuation phase and the date- T problem is the no-review problem of Section 3.2.

Two attainable date–menu pairs can be blended into a feasible random contract whenever one of them runs an incentive surplus and the other an incentive deficit. Write $b = (J, P)$ for the pair (10) generated by a date and its menu, and let b_+ and b_- be attainable pairs with $J_+ > 0 > J_-$. Mixing them with weight p on b_+ makes aggregate obedience bind at exactly one weight,

$$p(b_+, b_-) = \frac{-J_-}{J_+ - J_-} \in (0, 1), \quad (11)$$

at which the blend delivers the principal

$$\Pi(b_+, b_-) = p(b_+, b_-) P_+ + [1 - p(b_+, b_-)] P_- = \frac{-J_- P_+ + J_+ P_-}{J_+ - J_-}. \quad (12)$$

Note that p weights the *surplus* branch, whichever of the two dates that happens to be. The deterministic benchmark V^D is the yardstick for both statements below.

Proposition 8. *Suppose Assumptions 2 and 4 hold. Then $V^D \leq V^S$, and some optimal random contract uses at most two review dates. Moreover:*

(i) *randomization is strictly valuable, $V^S > V^D$, if and only if*

$$\Pi(b_+, b_-) > V^D \quad (13)$$

for some attainable pairs b_+, b_- with $J_+ > 0 > J_-$;

(ii) *whenever $V^S > V^D$, every minimal-support optimum is such a blend: it is supported on two distinct dates carrying attainable pairs b_+ and b_- with $J_+ > 0 > J_-$, mixed with the unique probability $p(b_+, b_-)$ of (11), and it attains $\Pi(b_+, b_-) = V^S$.*

Two dates suffice because there is one aggregate constraint. One date deliberately supplies more incentives than obedience requires and subsidizes another whose menu would not induce first-phase work on its own; only the average must be nonnegative. Parts (i) and (ii) apply the same blend to different objects: (i) ranges over arbitrary attainable pairs and merely certifies that randomization pays, whereas (ii) says the optimum is itself one such blend and collects the gain. In particular an optimal pair is always admissible in the test of (i), but a pair that passes the test need not be optimal.

Pointwise effort. Under the pointwise effort of Section 6.1 the scalar-constraint argument fails. If the earliest possible date passes without a review, the agent learns the remaining date and can revise his effort from then on, so a deviation must be evaluated allowing him to reoptimize. With two dates there is exactly one such event, and the problem stays tractable.

Fix $0 < \tau_1 < \tau_2 < T$ with probability p on τ_1 , put $d = \tau_2 - \tau_1$, and let r_i be the rent schedule used if review occurs at τ_i . Throughout this extension the two branches are indexed by *date*, not by the sign of the incentive surplus. If cumulative pre-review effort is e , write

$$U_i(e) = \mathbb{E}[r_i(e + \sqrt{\tau_i} Z)] - ce, \quad \widehat{U}_2(e) = \max_{m \in [e, e+d]} U_2(m), \quad Z \sim N(0, 1), \quad (14)$$

so that $\widehat{U}_2$ is the late-branch value *after* the agent has learned the date and reoptimized; the constraint $m \leq e + d$ is the flow bound on effort. Let

$$D_1(e) = U_1(\tau_1) - U_1(e), \quad D_2(e) = U_2(\tau_2) - \widehat{U}_2(e)$$

be the two branch losses from deviating to e before the first date.

Proposition 9. *Suppose Assumption 4 holds and effort is chosen pointwise, as in Section 6.1. Fix $0 < \tau_1 < \tau_2 < T$ and a weight p on τ_1 . Then:*

- (i) *Exact obedience. Full pre-review effort is sequentially incentive compatible if and only if $U_2(\tau_2) = \hat{U}_2(\tau_1)$ and*

$$pD_1(e) + (1-p)D_2(e) \geq 0 \quad \text{for every } e \in [0, \tau_1]. \quad (15)$$

- (ii) *A ceiling on the early weight. If in addition the late branch would induce full effort on its own, $U_2(\tau_2) \geq U_2(m)$ for every $m \in [0, \tau_2]$, then (15) holds if and only if $p \leq \bar{p}$, where*

$$\bar{p} = \inf_{\{e: D_1(e) < 0\}} \frac{D_2(e)}{D_2(e) - D_1(e)} \quad (16)$$

and $\bar{p} = 1$ when the set is empty.

- (iii) *Dominance. Let P_i be the branch payoffs (10) under full effort and let $V^{D, \text{pw}}$ be the best deterministic value among contracts that induce full pointwise effort. If $P_1 > V^{D, \text{pw}} > P_2$, a two-date lottery strictly dominates every such deterministic contract if and only if*

$$\underline{p} \equiv \frac{V^{D, \text{pw}} - P_2}{P_1 - P_2} < \bar{p}, \quad (17)$$

and any $p \in (\underline{p}, \bar{p}]$ achieves it.

The two bounds separate the two requirements. Weight $\underline{p}$ on the high-payoff branch is what beating one date costs; weight $\bar{p}$ is the most that branch can carry before some amount of early coasting becomes profitable. Randomization is valuable exactly when the payoff requirement is the weaker of the two. The phase-level problem checks fewer deviations, so $V^{D, \text{pw}} \leq V^D$ and V^D may be substituted for $V^{D, \text{pw}}$ in (17) to obtain a sufficient condition.

A numerical illustration. Both propositions can be illustrated with the same primitives, $T = 1.5$, $c = 0.72$ and $M = 2.60$, and with the same pair of review dates.

Under phase-level effort the best deterministic review occurs at $\tau^D \simeq 0.862$ and yields $V^D \simeq 0.07850$. The optimal lottery reviews at $\tau_1 \simeq 0.095$ and $\tau_2 \simeq 0.817$. The early date runs the incentive deficit and the later, more informative one supplies the offsetting surplus, so b_- sits at τ_1 and b_+ at τ_2 . By (11) the later date carries probability $p(b_+, b_-) \simeq 0.280$ and the early one the remaining 0.720, raising the payoff to $V^S \simeq 0.08279$ — about 5.5% above the deterministic benchmark.

Under pointwise effort, keep those two dates and let each menu assign A_B , A_D , $\tilde{A}_E$, A_G in order, with $\tilde{A}_E$ the pointwise easy action (9). Because the menus now use $\tilde{A}_E$ in place of A_E , the branch payoffs differ from those just reported: $P_1 = 0.13301$ at the early date and $P_2 = 0.03199$ at the late one. The late branch is itself pointwise incentive compatible, and $\bar{p} = 0.47293$. Since

$V^{D,pw} \leq V^D = 0.07850$, we get $\underline{p} \leq 0.46047 < \bar{p}$, so (17) holds; at $p = 0.465$ the lottery is worth 0.07896. A feasible two-date pointwise contract therefore beats every deterministic full-effort pointwise review, without having to solve the deterministic pointwise problem itself.

The improvement is small, and the second example is a witness rather than an optimized contract. Its point is qualitative: the value of random timing is not an artifact of phase-level effort. Even when the agent can coast before the first possible date, infers the late date from non-arrival, and then optimally revises his remaining effort, randomizing over two dates can strictly beat every deterministic review. Online Appendix OA.3 gives the proofs and the numerical construction.

7 Conclusion

This paper studies how a principal uses a public midterm review when effort is sequential and compensation is bounded. The review does more than evaluate past performance: it allows the principal to use future production, continuation rents, and termination as incentive instruments. The optimal menu assigns these continuation opportunities through an ordered performance ladder, while the review date balances the informativeness of accumulated performance against the incentive value of the remaining continuation phase.

Although three extensions demonstrate the robustness of our main findings, two further natural departures would enlarge the institution itself. First, in many organizations, the timing of the final evaluation is flexible: it may be brought forward before the deadline or postponed beyond the nominal deadline. Allowing the principal to choose a deadline at the midterm review would make the length of the continuation phase itself an incentive instrument. Performance realizations near the boundaries of the good- or bad-stop regions could then be assigned a short probationary period rather than being terminated immediately or retained until the deadline under a final performance standard. Second, our benchmark implicitly treats midterm reviews as sufficiently costly that at most one is conducted. Allowing multiple reviews would replace the four-rung menu with a backward recursion over convex continuation sets. Some faces of these sets would be inherited from the next review, while other regions would be strictly convexified, potentially yielding a rating rule that combines jumps, plateaus, and continuously graded bands. Both extensions enlarge the principal’s set of institutional instruments, and we leave their analysis for future work.

References

- Baker, G. P., M. C. Jensen, and K. J. Murphy (1988). Compensation and incentives: Practice vs. theory. *The Journal of Finance* 43(3), 593–616.
- Ball, I. and J. Knoepfle (2026). Should the timing of inspections be predictable? Revise and resubmit, Review of Economic Studies.
- Barlevy, G. and D. Neal (2019). Allocating effort and talent in professional labor markets. *Journal of Labor Economics* 37(1), 187–246.

- Cappelli, P. and M. J. Conyon (2018). What do performance appraisals do? *ILR Review* 71(1), 88–116.
- Chen, B. R. and Y. S. Chiu (2013). Interim performance evaluation in contract design. *The Economic Journal* 123(569), 665–698.
- Dai, L., Y. Wang, and M. Yang (2026). Dynamic contracting with flexible monitoring. Available at SSRN 3496785.
- Deserranno, E., P. Kastrau, and G. León-Ciliotta (2025). Promotions and productivity: The role of meritocracy and pay progression in the public sector. *American Economic Review: Insights* 7(1), 71–89.
- DeVaro, J. and D. Brookshire (2007). Promotions and incentives in nonprofit and for-profit organizations. *ILR Review* 60(3), 311–339.
- Ely, J. C., G. Georgiadis, and L. Rayo (2025). Feedback design in dynamic moral hazard. *Econometrica* 93(2), 597–621.
- Ely, J. C. and M. Szydlowski (2020). Moving the goalposts. *Journal of Political Economy* 128(2), 468–506.
- Georgiadis, G. and B. Szentes (2020). Optimal monitoring design. *Econometrica* 88(5), 2075–2107.
- Hoffmann, F., R. Inderst, and M. M. Opp (2021). Only time will tell: A theory of deferred compensation. *The Review of Economic Studies* 88(3), 1253–1278.
- Holmström, B. (1979). Moral hazard and observability. *The Bell Journal of Economics* 10(1), 74–91.
- Jewitt, I., O. Kadan, and J. M. Swinkels (2008). Moral hazard with bounded payments. *Journal of Economic Theory* 143(1), 59–82.
- Levin, J. and S. Tadelis (2005). Profit sharing and the role of professional partnerships. *The Quarterly Journal of Economics* 120(1), 131–171.
- Li, A. (2017). Efficiency in dynamic agency models. Working paper, Washington University in St. Louis.
- Liu, Y. (2026). Resource allocation with midterm review. Working paper.
- Lizzeri, A., M. A. Meyer, and N. Persico (2002). The incentive effects of interim performance evaluations. CARESS Working Paper 02-09, University of Pennsylvania.
- Lukas, C. (2010). Optimality of intertemporal aggregation in dynamic agency. *Journal of Management Accounting Research* 22(1), 157–174.

- Mas, A. (2017). Does transparency lead to pay compression? *Journal of Political Economy* 125(5), 1683–1721.
- Ray, K. (2007). Performance evaluations and efficient sorting. *Journal of Accounting Research* 45(4), 839–882.
- Rodivilov, A. (2022). Monitoring innovation. *Games and Economic Behavior* 135, 297–326.
- Sannikov, Y. (2008). A continuous-time version of the principal–agent problem. *The Review of Economic Studies* 75(3), 957–984.
- Spear, S. E. and C. Wang (2005). When to fire a CEO: Optimal termination in dynamic contracts. *Journal of Economic Theory* 120(2), 239–256.
- Varas, F., I. Marinovic, and A. Skrzypacz (2020). Random inspections and periodic reviews: Optimal dynamic monitoring. *The Review of Economic Studies* 87(6), 2893–2937.
- Wong, Y. F. (2026). Dynamic monitoring design. Revise and resubmit, American Economic Review.
- Zhu, J. Y. (2013). Optimal contracts with shirking. *The Review of Economic Studies* 80(2), 812–839.

Appendix

A Proofs for Section 3

Recall $\bar{\Delta}(h) = \max_z \Delta(h, z) = 2\Phi(\sqrt{h}/2) - 1$ and $M_I(h, c) = ch/\bar{\Delta}(h)$ from Section 3.1, and put $\bar{c}(h) = [2\Phi(\sqrt{h}/2) - 1]/\Phi(\sqrt{h}/2)$. Write

$$\ell_T = \frac{\phi(z_D(T) - \sqrt{T})}{\phi(z_D(T))} = \exp\left(z_D(T)\sqrt{T} - \frac{T}{2}\right), \quad \lambda_T^{\text{NR}} = \frac{1}{1 - \ell_T}, \quad (\text{A.1})$$

for the likelihood ratio at the no-review difficult bar and the shadow price of the no-review incentive constraint. Assumption 4 gives $z_D(T) < \sqrt{T}/2$, hence $\ell_T \in (0, 1)$ and $\lambda_T^{\text{NR}} > 1$.

Proof of Lemma 1. Let $k \in [-\infty, +\infty]$ solve $\mathbb{P}\{Y \geq k \mid e = 1\} = \mathbb{E}_1[w]/M$, which exists and is unique because that probability falls continuously and strictly from one to zero (with the endpoints understood as limits), and write $\varphi = w/M$, $\varphi^* = \mathbf{1}_{\{y \geq k\}}$. Since $\int(\varphi - \varphi^*)f_1^h = 0$, subtracting $\ell^h(k)$ times this vanishing integral gives

$$\mathbb{E}_0[w] - \mathbb{E}_0[\tilde{w}] = M \int (\varphi - \varphi^*)[\ell^h - \ell^h(k)]f_1^h.$$

For $y > k$, $\varphi^* = 1 \geq \varphi$ and $\ell^h(y) < \ell^h(k)$; for $y < k$ both inequalities reverse. The integrand is therefore nonnegative almost everywhere, so $\mathbb{E}_0[\tilde{w}] \leq \mathbb{E}_0[w]$, with equality only if $\varphi = \varphi^*$ a.e. Subtracting from $\mathbb{E}_1[\tilde{w}] = \mathbb{E}_1[w]$ gives the incentive comparison. $\square$

Proof of Lemma 2. Implementability. By Lemma 1 the largest attainable incentive spread is $M\bar{\Delta}(h)$, so $\mathcal{F}_h \neq \emptyset$ exactly when $M \geq M_I(h, c)$, and $[z_D, z_E]$ is nondegenerate exactly when the inequality is strict. With $a = \sqrt{h}/2$ we have $\bar{\Delta}(h)/h = \mathcal{G}(a) \equiv [2\Phi(a) - 1]/(4a^2)$ and $\mathcal{G}'(a) = -F(a)/(2a^3)$ with $F(a) = 2\Phi(a) - 1 - a\phi(a)$; since $F(0) = 0$ and $F'(a) = (1 + a^2)\phi(a) > 0$, $\mathcal{G}$ is strictly decreasing, so $M_I = c/\mathcal{G}$ is strictly increasing in h .

Profitability. By Lemma 1 and (1) the cheapest work-inducing contract is the difficult bar, so the principal's payoff is $V(M) = h - M\Phi(z_D)$. Dividing the binding condition $M[\Phi(z_D) - \Phi(z_D - \sqrt{h})] = ch$ by $\Phi(z_D)$,

$$M\Phi(z_D) = \frac{ch}{1 - q(z_D)}, \quad q(z) = \frac{\Phi(z - \sqrt{h})}{\Phi(z)}. \quad (\text{A.2})$$

Since ϕ/Φ is strictly decreasing, $\frac{d}{dz} \log q > 0$, so q is a strictly increasing bijection onto $(0, 1)$. As M rises, ch/M falls and z_D lies on the increasing branch of $\Delta(h, \cdot)$, so z_D falls, hence by (A.2) so does $M\Phi(z_D)$: the payoff V is strictly increasing in M , with $V \uparrow (1 - c)h > 0$ as $M \rightarrow \infty$ and, as $M \downarrow M_I(h, c)$ where the two bars merge at $\sqrt{h}/2$,

$$V \longrightarrow h \left[1 - \frac{c}{\bar{c}(h)}\right]. \quad (\text{A.3})$$

The threshold. If $c \leq \bar{c}(h)$ then (A.3) is nonnegative and $V > 0$ at every $M > M_I(h, c)$; set $\underline{M}(h, c) = M_I(h, c)$. If $c > \bar{c}(h)$ it is negative, so by strict monotonicity there is a unique $M_P(h, c) > M_I(h, c)$ with $V = 0$, and $\underline{M}(h, c) = M_P(h, c)$. At that cap $M\Phi(z_D) = h$ while the binding constraint gives $M\Phi(z_D - \sqrt{h}) = (1 - c)h$; dividing, z_D solves $q(z) = 1 - c$, and M_P is h over Φ of that root. In both cases $\underline{M} \geq M_I$, with equality exactly when $c \leq \bar{c}(h)$. $\square$

Proof of Proposition 1. By Lemma 1 and (1) the cheapest work-inducing contract is the difficult bar, whose expected cost under work is $M\Phi(z_D(T))$; thus, $V^{\text{NR}}(T; c, M) = T - M\Phi(z_D(T, c, M))$. Positivity follows from Lemma 2 under Assumption 4. The limits of V^{NR} and $z_D(T) \rightarrow -\infty$ are established in the proof of Lemma 2. $\square$

Proof of Lemma 3. Take $w \in \mathcal{F}_h$ and let $\tilde{w}$ be the cost-matched threshold contract of Lemma 1, which also lies in $\mathcal{F}_h$; its bar therefore satisfies (2), so $z \in [z_D, z_E]$ and $M\Phi(z_D) \leq \mathbb{E}_1[w] \leq M\Phi(z_E)$. Both bounds are attained, and $r = \mathbb{E}_1[w] - ch$ is affine in w on the convex set $\mathcal{F}_h$, so the attainable rents are exactly $[r_D, r_E]$, swept by mixtures of the two extreme contracts. The second equality is (1) at a binding constraint, and $r_E < M - ch$ because $\Phi(z_E) < 1$. Finally $\Delta(h, z) = \Phi(z) + \Phi(\sqrt{h} - z) - 1$ is symmetric about $\sqrt{h}/2$, so $z_E = \sqrt{h} - z_D$ and $r_E = M\Phi(-z_D) = M - M\Phi(z_D) = M - ch - r_D$. $\square$

The next lemma records two properties of the extreme continuation rents that are used repeatedly once the horizon is allowed to vary: they move continuously with the length of the continuation phase, and they spread to the full interval $[0, M]$ as that phase vanishes.

Lemma A.1. *Suppose Assumption 4 holds. Then $M\bar{\Delta}(h) > ch$ for every $h \in (0, T]$; the bars $z_D(h)$ and $z_E(h)$ of (3) are continuously differentiable in h there; and r_D and r_E are continuous on $(0, T]$ with*

$$\lim_{h \downarrow 0} r_D(h) = 0, \quad \lim_{h \downarrow 0} r_E(h) = M.$$

Consequently $h \mapsto \mathcal{A}(h)$ is continuous on $(0, T]$ in the Hausdorff metric, and $\mathcal{A}(h) \rightarrow \mathcal{A}(0) = \{(r, 0) : r \in [0, M]\}$ as $h \downarrow 0$.

Proof. Strict feasibility. Assumption 4 and Lemma 2 give $M > \underline{M}(T, c) \geq M_I(T, c)$, hence $M\bar{\Delta}(T) > cT$. Since $\bar{\Delta}(h)/h$ is strictly decreasing, as shown in the proof of Lemma 2, every $h \in (0, T)$ satisfies $\bar{\Delta}(h)/h > \bar{\Delta}(T)/T > c/M$. In particular the two roots are separated: $z_D(h) < \sqrt{h}/2 < z_E(h)$.

Continuity. Fix $h \in (0, T]$. Because $\Delta(h, \cdot)$ is strictly increasing to the left of $\sqrt{h}/2$ and strictly decreasing to its right, the partial derivative $\partial_z \Delta(h, z) = \phi(z) - \phi(z - \sqrt{h})$ is strictly positive at $z_D(h)$ and strictly negative at $z_E(h)$. Since Δ is C^1 in (h, z) , the implicit function theorem applied to $\Delta(h, z) = ch/M$ makes z_D and z_E continuously differentiable at h . By Lemma 3, $r_j(h) = M\Phi(z_j(h) - \sqrt{h})$ is then continuous.

Short-horizon limits. Let $h \downarrow 0$. Applying the mean value theorem to $\Delta(h, z) = \Phi(z) - \Phi(z - \sqrt{h})$ at the binding difficult bar yields $\xi_h \in (z_D(h) - \sqrt{h}, z_D(h))$ with $M\sqrt{h}\phi(\xi_h) = M\Delta(h, z_D(h)) = ch$, that is, $\phi(\xi_h) = c\sqrt{h}/M \rightarrow 0$, so $|\xi_h| \rightarrow \infty$. As $\xi_h < z_D(h) < \sqrt{h}/2 \rightarrow 0$, the family is eventually

bounded above, so $\xi_h \rightarrow -\infty$; and $z_D(h) < \xi_h + \sqrt{h}$ then gives $z_D(h) \rightarrow -\infty$. Hence $r_D(h) = M\Phi(z_D(h) - \sqrt{h}) \leq M\Phi(z_D(h)) \rightarrow 0$, and by Lemma 3, $r_E(h) = M - ch - r_D(h) \rightarrow M$.

The action set. $\mathcal{A}_0 = [0, M] \times \{0\}$ does not depend on h , while $\mathcal{A}_1(h) = [r_D(h), r_E(h)] \times \{(1-c)h\}$ is a segment whose endpoints and height all vary continuously in h ; Hausdorff continuity of the union follows. As $h \downarrow 0$ that segment converges to $[0, M] \times \{0\}$, which is $\mathcal{A}_0$. $\square$

B Proofs for Section 4

B.1 Review-score cutoffs and incentive supply

Section 4.2 establishes that, for a given $q = Q_\lambda(x)$, the maximizing action moves through A_B, A_D, A_E, A_G at the breakpoints $q_0 < q_1 = 0 < q_2$ in (6). For every $\lambda > 0$, Q_λ is strictly increasing in x and maps $\mathbb{R}$ onto $(-\infty, \lambda - 1)$. Hence the cutoff associated with q_j is finite exactly when $\lambda > 1 + q_j$, and inverting $Q_\lambda(k) = q_j$ gives

$$k_j(\lambda) = \frac{\tau}{2} - \log\left(1 - \frac{1 + q_j}{\lambda}\right) \quad \text{if } \lambda > 1 + q_j, \quad k_j(\lambda) = +\infty \quad \text{otherwise.} \quad (\text{B.1})$$

For later use, define $\underline{\lambda}(h) \equiv \max\{0, 1 + q_0\} = \max\{0, -V^{\text{NR}}(h)/r_D\}$, where $V^{\text{NR}}(h) = (1 - c)h - r_D$. Moreover, $1 + q_2 = (h + r_D)/(ch + r_D) < 1/c$, so any supporting multiplier $\lambda^* \geq 1/c$ makes all four regions active.

Define the standardized cutoff under first-phase work by $z_j(\lambda, \tau) \equiv (\tau - k_j(\lambda))/\sqrt{\tau}$. By (B.1),

$$z_j(\lambda, \tau) = \frac{\sqrt{\tau}}{2} + \frac{1}{\sqrt{\tau}} \log\left(1 - \frac{1 + q_j}{\lambda}\right), \quad j \in \{0, 1, 2\}, \quad (\text{B.2})$$

with $z_j = -\infty$ when $k_j = +\infty$. The induced rent schedule has jumps r_D , $r_E - r_D$, and $M - r_E$ at k_0, k_1, k_2 , so its first-phase incentive supply is

$$I(\lambda; \tau) = r_D \Delta(\tau, z_0) + [r_E - r_D] \Delta(\tau, z_1) + [M - r_E] \Delta(\tau, z_2), \quad (\text{B.3})$$

where $\Delta(\tau, -\infty) = 0$. Since each tie set $\{x : Q_\lambda(x) = q_j\}$ is null, (B.3) is unaffected by how ties are resolved.

B.2 The supporting multiplier

Lemma B.1. *Fix $\tau \in (0, T)$ and suppose Assumption 4 holds.*

- (i) *Monotonicity and range. $I(\lambda; \tau)$ is continuous and weakly increasing on $\lambda > 0$, and strictly increasing whenever the selected rent changes on a set of positive probability. Moreover, $\lim_{\lambda \rightarrow \infty} I(\lambda; \tau) = M\bar{\Delta}(\tau)$. If $(1 - c)h \neq r_D$, then also $\lim_{\lambda \downarrow \underline{\lambda}(h)} I(\lambda; \tau) = 0$.*
- (ii) *Existence. $M\bar{\Delta}(\tau) > c\tau$, and Problem (P_τ) admits a supporting multiplier $\lambda^* \geq 0$ satisfying complementary slackness.*

(iii) Positivity and binding obedience. If $(1-c)h \neq r_D$, every supporting multiplier satisfies $\lambda^* > 0$, and pre-review obedience binds:

$$\int_{\mathbb{R}} r^*(x) [g_1^\tau(x) - g_0^\tau(x)] dx = c\tau, \quad (\text{B.4})$$

where r^* is the rent delivered by the optimal menu. At the knife edge $(1-c)h = r_D$, a zero multiplier may also support an optimum.

(iv) Characterization. If $(1-c)h \neq r_D$, every supporting multiplier solves $I(\lambda; \tau) = c\tau$. The solution set is a nonempty compact interval contained in $(\underline{\lambda}(h), \infty)$, and is a singleton whenever $I(\cdot; \tau)$ is strictly increasing at that level.

Proof. (i) Let $\lambda'' > \lambda'$ and let (r', s') and (r'', s'') be Lagrangian maximizers at a given review score. Write $d_\tau(x) = 1 - g_0^\tau(x)/g_1^\tau(x)$. Adding the two optimality inequalities for $s - r + \lambda d_\tau r$ gives $(\lambda'' - \lambda')d_\tau(x)[r''(x) - r'(x)] \geq 0$. Multiplying by g_1^τ and integrating yields monotonicity of I , with strictness under the stated condition. Continuity follows from (B.2): when a cutoff first becomes finite, its standardized value tends to $-\infty$ and its contribution to (B.3) vanishes.

If $(1-c)h \neq r_D$, the limiting assignment becomes constant across review scores. When $1 + q_0 > 0$, it converges to A_B ; when $1 + q_0 < 0$, it converges to A_D . In either case $I \rightarrow 0$. Finally, as $\lambda \rightarrow \infty$, every $k_j \rightarrow \tau/2$, so every $z_j \rightarrow \sqrt{\tau}/2$; since the three rent jumps sum to M , $I \rightarrow M\bar{\Delta}(\tau)$.

(ii) By Lemma 2 and Assumption 4, $M > \underline{M}(T, c) \geq M_I(T, c) > M_I(\tau, c) = c\tau/\bar{\Delta}(\tau)$, where the strict inequality uses $\tau < T$ and monotonicity of $M_I(\cdot, c)$. Thus pre-review obedience is strictly feasible. Slater's condition and Kuhn–Tucker yield a supporting $\lambda^* \geq 0$ with complementary slackness.

(iii) Suppose $\lambda^* = 0$. Then $Q_0 \equiv -1$, so the review-score objective is $s - r$ and does not depend on x . If $(1-c)h \neq r_D$, its maximizer is unique: A_B when $(1-c)h < r_D$ and A_D when $(1-c)h > r_D$. The assigned rent is therefore constant across scores and supplies no first-phase incentives, contradicting $c\tau > 0$. Hence $\lambda^* > 0$, and complementary slackness gives (B.4).

At $(1-c)h = r_D$, instead, $q_0 = -1$, so under $\lambda = 0$ the bad stop A_B and difficult standard A_D are tied at every review score, while A_E and A_G are strictly dominated. A score-dependent selection between the tied actions can therefore provide the required first-phase incentives, so a zero supporting multiplier cannot be ruled out.

(iv) By (iii), $\lambda^* > 0$, so binding obedience is exactly $I(\lambda; \tau) = c\tau$. Parts (i)–(ii) imply that its solution set is nonempty, compact, and bounded away from $\underline{\lambda}(h)$. Weak monotonicity makes the level set an interval, and strict monotonicity at the level makes it a singleton. $\square$

B.3 The full ladder near the terminal date

The terminal boundary also yields a useful characterization of the shape of the fixed-review menu. The argument uses two limits. First, the threshold for activating the good stop converges to

one. Second, the fixed-review supporting multiplier converges to the shadow price of the terminal no-review problem, which is strictly larger than one.

Lemma B.2. *Suppose Assumption 4 holds and write $h = T - \tau$. Then, as $h \downarrow 0$,*

$$\frac{r_D(h)}{h} \longrightarrow +\infty, \quad q_2(h) \longrightarrow 0, \quad \lambda^*(T - h) \longrightarrow \lambda_T^{\text{NR}} > 1,$$

with ℓ_T and λ_T^{NR} as in (A.1).

Proof. For the first limit, the difficult continuation bar satisfies $M[\Phi(z_D(h)) - \Phi(z_D(h) - \sqrt{h})] = ch$. By the mean-value theorem there is $\xi_h \in (z_D(h) - \sqrt{h}, z_D(h))$ such that $M\sqrt{h}\phi(\xi_h) = ch$, that is, $\phi(\xi_h) = c\sqrt{h}/M$. On the difficult branch $\xi_h \rightarrow -\infty$, so this implies $|\xi_h| = O(\sqrt{\log(1/h)})$. Since $|z_D(h) - \sqrt{h} - \xi_h| \leq \sqrt{h}$, the corresponding normal densities are asymptotically equivalent. The lower Mills bound therefore gives, for some constant $C > 0$ and all sufficiently small h , $r_D(h) = M\Phi(z_D(h) - \sqrt{h}) \geq C\sqrt{h}/|\xi_h|$. Hence $r_D(h)/h \geq C/(\sqrt{h}|\xi_h|) \rightarrow +\infty$. Using $M - r_E(h) = ch + r_D(h)$, $q_2(h) = (1 - c)h/(ch + r_D(h)) \rightarrow 0$.

It remains to establish convergence of the supporting multiplier. Fix $1 < \lambda_- < \lambda_T^{\text{NR}} < \lambda_+$. By Lemma A.1, $r_D(h) \rightarrow 0$ and $r_E(h) \rightarrow M$ as $h \downarrow 0$. In the incentive-supply function (B.3), the first and third terms therefore vanish, while the middle term converges to $M\Delta(T, z_1(\lambda, T))$. At $\lambda = \lambda_T^{\text{NR}}$, (B.2) and (A.1) give $z_1(\lambda_T^{\text{NR}}, T) = z_D(T)$. The no-review incentive constraint therefore implies $M\Delta(T, z_1(\lambda_T^{\text{NR}}, T)) = cT$. By (A.1), $\ell_T \in (0, 1)$ and $\lambda_T^{\text{NR}} > 1$. Since $z_1(\cdot, T)$ is strictly increasing and $\Delta(T, \cdot)$ is strictly increasing to the left of $\sqrt{T}/2$, the limiting incentive supply is strictly increasing in λ at λ_T^{NR} . Hence, for all sufficiently small h , $I(\lambda_-; T - h) < c(T - h) < I(\lambda_+; T - h)$. Lemma B.1 then puts every supporting multiplier between λ_- and λ_+ . Since the bracket can be chosen arbitrarily tightly around λ_T^{NR} , $\lambda^*(T - h) \rightarrow \lambda_T^{\text{NR}}$. $\square$

B.4 Proofs

Proof of Proposition 2. By Lemma B.1, a supporting multiplier exists. Outside the knife edge it is strictly positive, so the review-score ranking derived in Section 4.2, together with the strict monotonicity of $Q_{\lambda^*}(x)$, gives the ladder structure of a^* : A_B if $x < k_0^*$, A_D if $k_0^* \leq x < k_1^*$, A_E if $k_1^* \leq x < k_2^*$, and A_G if $x \geq k_2^*$. The cutoff formula and the conditions for absent regions follow from (B.1). Since $q_0 < q_1 < q_2$, all finite cutoffs are strictly ordered.

At the knife edge $(1 - c)h = r_D$, the preceding argument applies whenever the optimum is supported by a positive multiplier. If instead $\lambda^* = 0$ supports an optimum, A_B and A_D are tied at every review score, while A_E and A_G are strictly dominated. Since first-phase work is implemented, an ordered threshold selection between A_B and A_D can be chosen to satisfy obedience. The proposition therefore continues to hold, in this case with $k_1^* = k_2^* = +\infty$. $\square$

Proof of Corollary 1. If $k_0^* = +\infty$, the assigned rent is identically zero and cannot induce first-phase work. Hence $k_0^* < +\infty$. Since Q_{λ^*} is continuous and strictly increasing, the difficult-standard region

is also nonempty. Finally, $Q_{\lambda^*}(\mathbb{R}) = (-\infty, \lambda^* - 1)$ reaches $q_1 = 0$ iff $\lambda^* > 1$ and reaches q_2 iff $\lambda^* > 1 + q_2$, which gives the three cases in the corollary. $\square$

Proof of Proposition 3. Write $h = T - \tau$. By Lemma B.2, $\lambda^*(T - h) \rightarrow \lambda_T^{\text{NR}} > 1$ and $1 + q_2(h) \rightarrow 1$. Hence, for all sufficiently small $h > 0$, $\lambda^*(T - h) > 1 + q_2(h)$. Moreover, $r_D(h)/h \rightarrow \infty$ implies $r_D(h) \neq (1 - c)h$ for all sufficiently small h , so the generic cutoff characterization in Corollary 1 applies. Therefore $k_2^*(T - h) < +\infty$ for every sufficiently small $h > 0$, as claimed. $\square$

C Proofs for Section 5

C.1 The fixed-review value and the one-sided derivatives

Let $u_D(h) = (1 - c)h - r_D(h)$ and $u_E(h) = (1 - c)h - r_E(h)$ be the principal's payoffs under the two continuation standards, with the argument suppressed where the horizon is fixed. Conditional on first-phase work, the probabilities of the four regions are determined by the standardized cutoffs, and the principal's value from the optimal fixed- τ menu is

$$\begin{aligned} V^R(\tau) = & \tau + u_D [\Phi(z_0(\lambda, \tau)) - \Phi(z_1(\lambda, \tau))] \\ & + u_E [\Phi(z_1(\lambda, \tau)) - \Phi(z_2(\lambda, \tau))] - M\Phi(z_2(\lambda, \tau)), \end{aligned} \quad (\text{C.1})$$

where $\lambda = \lambda^*(\tau)$ is a supporting multiplier and $\Phi(-\infty) = 0$. Equation (C.1) applies to the two-, three-, and four-region cases under the conventions of (B.2).

For interior review dates, continuity follows from the implicit-function characterization of $z_D(h)$ and $z_E(h)$, continuity of the review-action set, and the maximum theorem applied to the fixed-review linear program. Changes in the number of review regions may create changes in the formula used to compute the value, but not jumps in the value itself.

By Lemma 4 the review gain $G(\tau) = V^R(\tau) - V^{\text{NR}}(T; c, M)$ vanishes at both boundaries, and the natural objects to study there are the one-sided derivatives

$$V^{R'}(0_+) := \lim_{\tau \downarrow 0} \frac{V^R(\tau) - V^{\text{NR}}(T; c, M)}{\tau}, \quad V_-^{R'}(T) := \lim_{\tau \uparrow T} \frac{V^R(\tau) - V^{\text{NR}}(T; c, M)}{\tau - T}, \quad (\text{C.2})$$

whenever the corresponding limits exist. These definitions do not extend V^R to either boundary.

C.2 The terminal derivative

Propositions 4(ii) and 6 turn on the sign of $V_-^{R'}(T)$, computed here. Throughout, $z_D(T)$ is the difficult bar of (3) at horizon T , and ℓ_T and λ_T^{NR} are defined in (A.1).

Lemma C.1. *Suppose Assumption 4 holds. Then $V_-^{R'}(T)$ exists, the review value and the no-review value agree to first order at T ,*

$$V^R(T - h) = V^{\text{NR}}(T - h; c, M) + o(h), \quad V_-^{R'}(T) = \frac{\partial V^{\text{NR}}(T; c, M)}{\partial T},$$

and

$$V_-^{R'}(T) = 1 - \lambda_T^{\text{NR}} \left[c - \frac{M\phi(z_D(T) - \sqrt{T})}{2\sqrt{T}} \right]. \quad (\text{C.3})$$

Consequently:

(i) $V_-^{R'}(T) < 0$ if and only if

$$\ell_T - (1 - c) > \frac{M\phi(z_D(T) - \sqrt{T})}{2\sqrt{T}}; \quad (\text{C.4})$$

(ii) $V_-^{R'}(T) < 0$ implies $\lambda_T^{\text{NR}} > 1/c$, equivalently $\ell_T > 1 - c$.

Proof. As $h \downarrow 0$, the work-inducing continuation frontier differs from the no-effort frontier by a vertical surplus of only $(1 - c)h$. Moreover, r_D and $M - r_E$ are both $O(\sqrt{h}/\sqrt{\log(1/h)})$. The continuation actions can therefore improve on the no-effort support function only on likelihood-index intervals whose widths are $O(\sqrt{h \log(1/h)})$. Since the density of the review signal is bounded, their total incremental value is $o(h)$, which gives $V^R(T - h) = V^{\text{NR}}(T - h; c, M) + o(h)$. With $\lim_{\tau \uparrow T} V^R(\tau) = V^{\text{NR}}(T; c, M)$ from Lemma 4 and definition (C.2), this yields $V_-^{R'}(T) = \partial V^{\text{NR}}(T; c, M)/\partial T$.

It remains to differentiate the no-review value. Implicit differentiation of $M[\Phi(z_D(T)) - \Phi(z_D(T) - \sqrt{T})] = cT$ gives

$$\frac{dz_D(T)}{dT} = \frac{c/M - \phi(z_D(T) - \sqrt{T})/(2\sqrt{T})}{\phi(z_D(T)) - \phi(z_D(T) - \sqrt{T})}.$$

Differentiating $V^{\text{NR}}(T; c, M) = T - M\Phi(z_D(T))$ and using (A.1) gives (C.3).

For (i), substitute $\lambda_T^{\text{NR}} = 1/(1 - \ell_T)$ into (C.3) and multiply by $1 - \ell_T > 0$: $V_-^{R'}(T) < 0$ is equivalent to $c - M\phi(z_D(T) - \sqrt{T})/(2\sqrt{T}) > 1 - \ell_T$, which is (C.4). For (ii), the term $\lambda_T^{\text{NR}} M\phi(z_D(T) - \sqrt{T})/(2\sqrt{T})$ in (C.3) is strictly positive, so $V_-^{R'}(T) < 0$ forces $1 - \lambda_T^{\text{NR}} c < 0$; and $\lambda_T^{\text{NR}} > 1/c$ is equivalent to $1 - \ell_T < c$. $\square$

The left side of (C.4) is large when the reward event is weak evidence of work, that is when the shadow price of terminal incentives is high; the right side is the first-order informational gain from lengthening the horizon. Part (ii) is (C.4) with that right side replaced by zero, and hence weaker. Recall from Lemma 2 that $\underline{M}(T, c) = M_I(T, c)$ when $c \leq \bar{c}(T)$.

Lemma C.2. Fix $T > 0$ and $c \in (0, 1)$. Then

$$\lim_{M \downarrow M_I(T, c)} V_-^{R'}(T) = -\infty, \quad \lim_{M \rightarrow \infty} V_-^{R'}(T) = 1 - c > 0. \quad (\text{C.5})$$

Proof. Let $a = \sqrt{T}/2$. As $M \downarrow M_I(T, c)$ the two no-review roots merge, so $z_D(T) \uparrow a$, $\ell_T \uparrow 1$, and $\lambda_T^{\text{NR}} \rightarrow \infty$. The bracket in (C.3) converges to $c[1 - a\phi(a)/(2\Phi(a) - 1)]$, which is strictly positive because

$$2\Phi(a) - 1 - a\phi(a) > 0 \quad \text{for every } a > 0. \quad (\text{C.6})$$

Indeed, the expression on the left of (C.6) is zero at $a = 0$ and has derivative $(1 + a^2)\phi(a) > 0$. The first limit in (C.5) follows.

As $M \rightarrow \infty$, $z_D(T) \rightarrow -\infty$ by the proof of Lemma 2; hence $\ell_T \rightarrow 0$, $\lambda_T^{\text{NR}} \rightarrow 1$, and $M\phi(z_D(T) - \sqrt{T}) \rightarrow 0$. Equation (C.3) then gives the second limit. $\square$

Corollary C.1. *Suppose $c \leq \bar{c}(T)$, so that $\underline{M}(T, c) = M_I(T, c)$. Then there exists $\delta > 0$ such that $V_-^{Rl}(T) < 0$ for every $M \in (\underline{M}(T, c), \underline{M}(T, c) + \delta)$.*

Proof. Immediate from the first limit in (C.5). $\square$

C.3 Proofs

Proof of Lemma 4. Consider first $\tau \downarrow 0$. The largest principal payoff available from a second-phase action is $u_D = (1 - c)h - r_D$ for all sufficiently small τ , because profitability of the no-review contract implies $u_D(T) = V^{\text{NR}}(T; c, M) > 0$. Hence $V^R(\tau) \leq \tau + u_D(T - \tau)$, whose limit is $u_D(T) = V^{\text{NR}}(T; c, M)$.

For a matching lower bound, assign the low-rent continuation action A_D at every review score except on a sufficiently high upper tail, on which the principal assigns the good stop A_G . The utility increment on that tail is $M - r_D$. Choose the upper-tail cutoff so that $[M - r_D]\Delta(\tau, z) = c\tau$. For sufficiently small τ , such a cutoff exists, and the work probability of the reward region converges to zero. The resulting payoff converges to $u_D(T)$. Therefore $\lim_{\tau \downarrow 0} V^R(\tau) = V^{\text{NR}}(T; c, M)$.

Now let $\tau \uparrow T$, so $h \downarrow 0$. A menu using only A_B and A_G is exactly a bounded upper-tail bonus based on X_τ . Choosing its threshold to satisfy $M\Delta(\tau, z) = c\tau$ gives payoff $V^{\text{NR}}(\tau; c, M)$, which converges to $V^{\text{NR}}(T; c, M)$.

For the reverse inequality, the full-effort continuation frontier converges to the no-effort frontier. In particular $(1 - c)h \rightarrow 0$ and, by Lemma A.1, $r_D \rightarrow 0$ and $r_E \rightarrow M$. Thus the complete review-action set converges to $\{(r, 0) : r \in [0, M]\}$, so the associated principal payoff converges to $-r$. The review signal converges to the terminal signal $N(T, T)$, and the fixed-review problem converges to the no-review bounded-payment problem. This proves $\lim_{\tau \uparrow T} V^R(\tau) = V^{\text{NR}}(T; c, M)$. $\square$

Proof of Proposition 4. The two claims concern the one-sided derivatives (C.2). Part (i) establishes that $V^{Rl}(0_+) = -\infty$ directly; part (ii) reads the sign of $V_-^{Rl}(T)$ off Lemma C.1.

For the first claim, write $h = T - \tau$ and use the low-rent continuation action $A_D(h)$ as a baseline. For all sufficiently small τ , this action uniquely maximizes the principal's post-review payoff. Compactness of the two continuation frontiers and strict profitability at $h = T$ imply that there is a constant $\kappa > 0$ such that every feasible review action (r, s) satisfies

$$u_D - (s - r) \geq \kappa|r - r_D|. \quad (\text{C.7})$$

Let $L_\tau(x) = 1 - g_0^\tau(x)/g_1^\tau(x)$. Because a constant review utility provides no incentives, pre-review obedience can be written as

$$\mathbb{E}_1 [(r(X_\tau) - r_D)L_\tau(X_\tau)] \geq c\tau. \quad (\text{C.8})$$

Under work, $X_\tau = \tau + \sqrt{\tau}Z$ and $L_\tau(X_\tau) = 1 - \exp(-\tau/2 - \sqrt{\tau}Z)$. A standard Gaussian tail decomposition gives

$$\mathbb{E}_1[|r(X_\tau) - r_D|] \geq K \frac{\sqrt{\tau}}{\sqrt{\log(1/\tau)}} \quad (\text{C.9})$$

for some $K > 0$ and all sufficiently small τ . To see the order, split at $|Z| = \sqrt{8 \log(1/\tau)}$. On the central event, $|L_\tau|$ is of order $\sqrt{\tau \log(1/\tau)}$; under both work and shirking, the complementary Gaussian tails contribute $o(\tau)$. Since review utilities are bounded by M , (C.8) then implies (C.9).

Combining (C.7) and (C.9) yields

$$V^R(\tau) \leq \tau + u_D(T - \tau) - \kappa K \frac{\sqrt{\tau}}{\sqrt{\log(1/\tau)}}.$$

Since $\tau + u_D(T - \tau) = V^{\text{NR}}(T; c, M) + O(\tau)$, it follows that

$$\lim_{\tau \downarrow 0} \frac{V^R(\tau) - V^{\text{NR}}(T; c, M)}{\tau} = -\infty. \quad (\text{C.10})$$

Thus $V^R(\tau) < V^{\text{NR}}(T; c, M)$ for all sufficiently small $\tau > 0$. Since the boundary limit established above gives $V^R(s) \rightarrow V^{\text{NR}}(T; c, M)$ as $s \downarrow 0$, each such τ is strictly dominated by some genuine review date $s \in (0, \tau)$. This proves part (i).

For the second claim, let M lie in the right neighborhood supplied by Corollary C.1, so that $V_-^{R'}(T) < 0$. By Lemma C.1, $V^R(T - h) = V^{\text{NR}}(T; c, M) - V_-^{R'}(T)h + o(h)$. Comparing the dates $\tau = T - h$ and $\tau' = T - 2h$, $V^R(\tau') - V^R(\tau) = -V_-^{R'}(T)h + o(h) > 0$ for all sufficiently small $h > 0$. Hence every sufficiently late review is strictly dominated by an earlier genuine review, which proves the second claim. $\square$

Proof of Proposition 5. Write $R_M(t) \equiv R^{\text{NR}}(t; c, M) = r_D(t)$. By Lemma 2, $R_M(T) \rightarrow 0$ as $M \rightarrow \infty$, so the no-review full-effort contract is feasible and profitable for all sufficiently large M .

Suppose toward a contradiction that there are sequences $M_n \rightarrow \infty$ and $\tau_n \in (0, T)$ such that

$$V^R(\tau_n) \geq V^{\text{NR}}(T; c, M_n). \quad (\text{C.11})$$

Dependence on n is suppressed below. Any continuation-utility schedule $r(X_\tau) \in [0, M]$ that induces first-phase work is itself a bounded incentive contract for a phase of length τ . By the definition of $R_M(\tau)$, the agent's ex ante rent under the review contract is therefore at least $R_M(\tau)$. Total surplus is at most $(1 - c)T$, and hence $V^R(\tau) \leq (1 - c)T - R_M(\tau)$. Comparison with (C.11) yields the necessary condition

$$R_M(\tau) \leq R_M(T). \quad (\text{C.12})$$

We first show that any such sequence must satisfy

$$\tau_n \log M_n \rightarrow 0. \quad (\text{C.13})$$

Implicit differentiation of the binding difficult-root condition gives

$$R'_M(t) = \frac{\phi(z_D(t) - \sqrt{t})}{\phi(z_D(t)) - \phi(z_D(t) - \sqrt{t})} \left[c - \frac{M\phi(z_D(t))}{2\sqrt{t}} \right]. \quad (\text{C.14})$$

For every $\varepsilon > 0$, $z_D(t) \rightarrow -\infty$ uniformly on $t \in [\varepsilon, T]$ as $M \rightarrow \infty$. The Mills bound $\Phi(z) \leq \phi(z)/|z|$ for $z < 0$, together with $M\Delta(t, z_D(t)) = ct$, makes the bracket in (C.14) negative uniformly on that interval. The prefactor is positive on the difficult branch. Hence $R'_M(t) < 0$ on $[\varepsilon, T]$ for all sufficiently large M . Condition (C.12) then forces $\tau_n \rightarrow 0$.

We sharpen this localization. Let $a_M < 0$ solve $\Phi(a_M) = 2cT/M$, and put $q_M = \Phi(a_M - \sqrt{T})/\Phi(a_M)$. The threshold bonus $M/[2(1 - q_M)]$ is feasible for large M , implements effort, and leaves rent $cTq_M/(1 - q_M)$. Gaussian tail bounds therefore give

$$R_M(T) \leq 2cT \exp \left(-\sqrt{T \log M} - \frac{T}{2} \right) \quad (\text{C.15})$$

for all sufficiently large M .

For a short phase of length τ , write $z_D(\tau) = -x$ and $u = \sqrt{\tau}$. The binding condition implies $R_M(\tau) = c\tau q/(1 - q) \geq c\tau q$, where $q = \Phi(-(x + u))/\Phi(-x)$. The two-sided Mills inequalities and $x \leq \sqrt{2 \log(M/(c\tau))}$ give

$$R_M(\tau) \geq \frac{c\tau}{2(1 + \sqrt{T})} \exp \left[-\sqrt{2\tau \log \left(\frac{M}{c\tau} \right)} - \frac{\tau}{2} \right]. \quad (\text{C.16})$$

If $\tau_n \log M_n$ were bounded away from zero along a subsequence, then the logarithm of the right side of (C.16) would be $-o(\sqrt{\log M_n})$, whereas (C.15) is of order $-\sqrt{T \log M_n}$. It would follow that $R_{M_n}(\tau_n) > R_{M_n}(T)$, contradicting (C.12). This proves (C.13).

It remains to show that a review in this vanishingly early region is also strictly dominated. Put $h = T - \tau$, $r_0 = R_M(h)$, and let the difficult continuation action be (r_0, s_0) , where $s_0 = (1 - c)h$ is total surplus and its principal payoff is $s_0 - r_0$. Uniformly for $h \in [T/2, T]$, $r_0 \rightarrow 0$; hence $2r_0 < (1 - c)h$ for all sufficiently large M . Every continuation action $(r, s) \in \mathcal{A}(h)$ then satisfies

$$s - r \leq s_0 - r_0 - |r - r_0|. \quad (\text{C.17})$$

For a work-inducing action $s = (1 - c)h = s_0$ and $r \geq r_0$, so the inequality holds with equality. For a no-effort action $s = 0$; it follows immediately when $r \geq r_0$, and when $r < r_0$ it follows from $2r_0 < (1 - c)h$.

Let $d(x) = r(x) - r_0$. Integrating (C.17) under first-phase work gives

$$V^R(\tau) \leq \tau + s_0 - r_0 - \mathbb{E}_1 |d(X_\tau)|. \quad (\text{C.18})$$

Since a constant rent contributes nothing to incentives, first-phase obedience is

$$\mathbb{E}_1[d(X_\tau)L_\tau(X_\tau)] \geq c\tau, \quad (\text{C.19})$$

Under work, $X_\tau = \tau + \sqrt{\tau}Z$ and $L_\tau(X_\tau) = 1 - \exp(-\tau/2 - \sqrt{\tau}Z)$. On the event $|L_\tau| \leq 1/2$, the left side of (C.19) is bounded by $\mathbb{E}_1|d|/2$. On its complement, $|d| \leq M$, and Gaussian tail bounds give, with $a = \log(3/2)$, $\mathbb{E}_1[|L_\tau|\mathbf{1}_{\{|L_\tau| > 1/2\}}] \leq 4\exp(-a^2/(8\tau))$. Because (C.13) is equivalent to $\log M = o(1/\tau)$, this tail contribution is $o(\tau)$. Thus

$$\mathbb{E}_1|d(X_\tau)| \geq 2c\tau - o(\tau) > c\tau \quad (\text{C.20})$$

for all sufficiently large n .

Finally, subtracting $V^{\text{NR}}(T; c, M) = (1 - c)T - R_M(T)$ from (C.18) yields

$$V^R(\tau) - V^{\text{NR}}(T; c, M) \leq c\tau + R_M(T) - R_M(T - \tau) - \mathbb{E}_1|d(X_\tau)|.$$

For large n , $T - \tau_n \in [T/2, T)$ and the monotonicity established after (C.14) gives $R_M(T - \tau) > R_M(T)$. Combining this strict inequality with (C.20) makes the displayed difference negative, contradicting (C.11). No such sequence exists, so a finite uniform $\overline{M}(T, c)$ exists. $\square$

Proof of Proposition 6. By Lemma 4, $G(\tau) \rightarrow 0$ as $\tau \uparrow T$, and $V_-^{Rl}(T)$ in (C.2) is exactly the left derivative of G at T :

$$G(\tau) = V_-^{Rl}(T)(\tau - T) + o(T - \tau) \quad \text{as } \tau \uparrow T.$$

Under the stated hypotheses, Corollary C.1 gives $V_-^{Rl}(T) < 0$. Since $\tau - T < 0$, the leading term is strictly positive, so there is $\delta > 0$ with $G(\tau) > 0$ for every $\tau \in (T - \delta, T)$, which is the claim. $\square$

The same display gives the converse: if $V_-^{Rl}(T) > 0$ then $G(\tau) < 0$ throughout a left neighborhood of T .

Proof of Proposition 7. At the boundary, $\overline{\Delta}(T)/T = c/M_0$. Since $\overline{\Delta}(t)/t$ is strictly decreasing by Lemma 2, every $t < T$ satisfies $\overline{\Delta}(t)/t > c/M_0$, which gives $M_0\overline{\Delta}(\tau) > c\tau$ and $M_0\overline{\Delta}(h) > ch$.

To obtain the second claim, choose $\tau_0 > 0$ small and write $h_0 = T - \tau_0$. At $M = M_0$ the continuation problem at $h_0 < T$ is strictly feasible, so its minimum work-inducing rent $r_D(h_0)$ is well defined and strictly below M_0 . At the boundary $h = T$, the two threshold roots coincide and the corresponding rent is $M_0\Phi(-\sqrt{T}/2)$; hence the available rent spread converges to $M_0\Phi(\sqrt{T}/2) > 0$ as $h_0 \uparrow T$. Since $\overline{\Delta}(\tau)/\tau \rightarrow \infty$ as $\tau \downarrow 0$, τ_0 can therefore be chosen small enough that

$$[M_0 - r_D(h_0)]\overline{\Delta}(\tau_0) > c\tau_0. \quad (\text{C.21})$$

Both continuation feasibility at h_0 and (C.21) are strict, so by continuity they continue to hold for all caps in a sufficiently small interval $(M_0 - \varepsilon, M_0]$.

For such a cap, use the minimum-rent full-effort continuation contract after review scores $X_{\tau_0} < \tau_0/2$, and use the good-stop action with rent M after $X_{\tau_0} \geq \tau_0/2$. The spread in continuation utility between the two regions is $M - r_D(h_0)$, and the midpoint cutoff maximizes the difference between the work and shirk probabilities. The date-0 incentive spread is therefore $[M - r_D(h_0)]\bar{\Delta}(\tau_0) > c\tau_0$, so first-phase work is induced. In the lower region the continuation contract induces full effort, while in the upper region the project is stopped. The lower region has strictly positive probability under work. Finally, $M < M_0 = M_I(T, c)$ makes full effort infeasible without review by the implementability condition in Lemma 2. $\square$

Online Appendix for “Midterm Review”

Doruk Cetemen Yonggyun Kim Fei Li Curtis R. Taylor

OA.1 Pointwise effort choice

Section 6.1 states the two conclusions of this extension: no interior effort level becomes an exposed continuation action, and the easy continuation standard must be raised from $z_E(h)$ to $\zeta(h)$, lowering the maximum continuation rent from $r_E(h)$ to $\tilde{r}_E(h)$. This appendix supplies the argument.

OA.1.1 Flexible effort and a relaxed feasible set

Fix a continuation phase of length h and let total effort be $\eta \in [0, h]$, so that a plan with total effort η produces $Y \sim N(\eta, h)$ and costs the agent $c\eta$. For a terminal wage w write $\pi(\eta) = \mathbb{E}_\eta[w] = \int_{\mathbb{R}} w(y) f_\eta(y) dy$ for the agent’s expected payment when he supplies total effort η , and let $\mathcal{F}_h^{\text{flex}}(\eta)$ denote the contracts satisfying $(\text{IC}_h^{\text{flex}})$. Proofs are collected in Section OA.1.6.

The economic difference from phase-level effort is that the old model asks only whether the agent wants to work rather than not work, whereas the flexible model also asks whether he wants to work *all the way to the target*. Those nearby deviations are what tighten the easy standard below.

A relaxed feasible set. Condition $(\text{IC}_h^{\text{flex}})$ contains a continuum of deviations, and we will not need to characterize all of them. Three necessary conditions suffice. Two are already in the main text: (G_η) , the comparison with abandoning effort altogether, which is exactly the comparison used in the phase-level model; and (L_η) , which, because $\pi'(\eta) = \text{Cov}_\eta(w(Y), Y)/h$, says that the wage must load enough on output to pay for the *marginal* unit of effort. The second is the new restriction that matters for the easy standard. Third, at an interior target the agent’s payoff must bend the right way:

$$\pi''(\eta) \leq 0 \quad \text{if } \eta \in (0, h). \quad (\text{S}_\eta)$$

Let $\mathcal{F}_h(\eta)$ impose only (G_η) , and let $\hat{\mathcal{F}}_h(\eta)$ impose the global, local, and second-order conditions. Then

$$\mathcal{F}_h^{\text{flex}}(\eta) \subseteq \hat{\mathcal{F}}_h(\eta) \subseteq \mathcal{F}_h(\eta). \quad (\text{OA.1})$$

The rightmost set is the old binary-effort benchmark. The middle set is an analytically useful outer bound on the true flexible-effort set. The strategy below is to show that even this larger set has no economically relevant extreme points beyond the four actions we identify. Since the four candidate vertices themselves satisfy the full incentive constraint, this sandwich argument is enough; we never have to solve the entire continuum-deviation problem at interior effort.

Rent and surplus. As in the baseline model, an action is the pair (r, s) of continuation rent and *total* continuation surplus, the principal's own payoff being $u = s - r$. If the contract implements η ,

$$r = \pi(\eta) - c\eta, \quad s = (1 - c)\eta, \quad u = s - r = \eta - \pi(\eta). \quad (\text{OA.2})$$

For a fixed effort target, feasible contracts therefore lie on a horizontal segment at height $(1 - c)\eta$, along which the principal's payoff falls one for one with the rent conceded. Flexible effort fills in all the segments corresponding to $\eta \in [0, h]$:

$$\mathcal{A}^{\text{flex}}(h) = \bigcup_{\eta \in [0, h]} \{(P - c\eta, (1 - c)\eta) : P \in [P_{\min}(\eta), P_{\max}(\eta)]\}, \quad (\text{OA.3})$$

where $P_{\min}(\eta)$ and $P_{\max}(\eta)$ are the least and greatest expected payments that can implement η . Define $\hat{P}_{\min}$ and $\hat{P}_{\max}$ analogously using the relaxed set $\hat{\mathcal{F}}_h(\eta)$.

The picture is therefore richer than in the baseline model: instead of two line segments, the principal faces a two-dimensional body. The key question, however, is not how large that body is. It is whether the body creates new *extreme* rent-surplus pairs. Since the review problem maximizes a linear function of (r, s) state by state, only those extreme points can ever matter for the optimal menu.

OA.1.2 The three regimes

There are three economically distinct targets: partial effort, full effort, and no effort. Partial effort is the new case. Full effort turns out to look almost exactly like the baseline model, while no effort is unchanged.

Interior effort: why a band replaces a threshold. Suppose the principal wants to implement $\eta \in (0, h)$. In the baseline model an upper-tail bonus is attractive because high output is evidence of effort. With an interior target that logic is incomplete: a very high outcome is also evidence that the agent supplied *more* than the target. A contract that keeps paying on arbitrarily high outcomes therefore encourages the wrong upward deviation.

The least-cost way to induce an interior target consequently rewards a *range* of good outcomes rather than the entire upper tail. The lower edge of the range makes effort attractive; the upper edge prevents the agent from wanting to overshoot the target. Formally, matching both the expected payment and the marginal incentive produces the following band-reduction result.

Lemma OA.1. *Fix $\eta \in (0, h)$ and let $w : \mathbb{R} \rightarrow [0, M]$ satisfy $\pi'_w(\eta) = c$. Then there is a band bonus*

$$\tilde{w}(y) = M \mathbf{1}\{k_1 \leq y \leq k_2\}, \quad -\infty \leq k_1 < k_2 \leq +\infty, \quad (\text{OA.4})$$

with the same expected payment and the same marginal incentive at the target,

$$\pi_{\tilde{w}}(\eta) = \pi_w(\eta), \quad \pi'_{\tilde{w}}(\eta) = c, \quad \text{and} \quad \pi_{\tilde{w}}(0) \leq \pi_w(0).$$

Consequently $w \in \hat{\mathcal{F}}_h(\eta)$ implies $\tilde{w} \in \hat{\mathcal{F}}_h(\eta)$, so the range of $\pi(\eta)$ over $\hat{\mathcal{F}}_h(\eta)$ is unchanged if attention is restricted to band bonuses.

The economic content is simple. The principal still wants to spend the payment cap where output is informative about effort. But she now has to control incentives on both sides of the target. The optimal payment region is therefore bounded above as well as below: outcomes that are sufficiently good earn the bonus, while extremely good outcomes do not. The upper cutoff is not a punishment for success; it is what prevents the contract from rewarding effort beyond the amount the principal wants to buy.

Lemma OA.1 reduces the problem to two numbers. Standardize the band by

$$a = \frac{k_1 - \eta}{\sqrt{h}}, \quad b = \frac{k_2 - \eta}{\sqrt{h}}, \quad \kappa = \frac{c\sqrt{h}}{M},$$

so that $\pi(\eta) = M[\Phi(b) - \Phi(a)]$ and $\pi'(\eta) = M[\phi(a) - \phi(b)]/\sqrt{h}$. The interior condition in (L_η) then reads

$$\phi(a) - \phi(b) = \kappa, \tag{OA.5}$$

a single equation in (a, b) that does not involve η . Let $\zeta(h) > 0$ be the bar of Section 6.1, the positive root of $\phi(\zeta) = \kappa$, equivalently of $M\phi(\zeta)/\sqrt{h} = c$. Along the curve (OA.5) the payment $M[\Phi(b) - \Phi(a)]$ is largest at $a = -\zeta(h)$, $b = +\infty$: the band degenerates to an upper-tail bonus with bar $\eta - \zeta(h)\sqrt{h}$, and the payment it delivers,

$$p^* = M\Phi(\zeta(h)), \tag{OA.6}$$

uses (L_η) alone and does not vary with η . The bound is attained, and the slice is an interval.

Lemma OA.2. *Suppose $\Delta(h, \zeta(h)) \geq ch/M$. Then for every $\eta \in (0, h)$ the expected payments achievable at η form the interval $[P_{\min}(\eta), p^*]$. The upper end is attained by the upper-tail bonus with bar $\eta - \zeta(h)\sqrt{h}$, which lies in $\mathcal{F}_h^{\text{flex}}(\eta)$ itself and not merely in the relaxation. The lower end is $P_{\min}(\eta) = M[\Phi(b) - \Phi(a)]$ at the solution of (OA.5) together with (G_η) at equality: writing $\theta = \eta/\sqrt{h}$,*

$$[\Phi(b) - \Phi(a)] - [\Phi(b + \theta) - \Phi(a + \theta)] = \kappa\theta. \tag{OA.7}$$

Because the upper end does not vary with η , the upper boundary of $\mathcal{A}^{\text{flex}}(h)$ is a straight line. Because the maximum expected payment p^* is the same for every interior target, the high-rent side of the feasible set has a particularly simple interpretation. Along that side, the principal is already paying as much as the local incentive condition permits. Increasing the effort target therefore means asking the agent to work more without increasing expected compensation. Substituting $P = p^*$ into (OA.2) gives $(r, s) = (p^* - c\eta, (1 - c)\eta)$ for $\eta \in [0, h]$, which runs from the no-effort face to

$$\tilde{A}_E = (p^* - ch, (1 - c)h) \tag{OA.8}$$

at full effort.

This already suggests why partial effort will not survive in the final menu. Every interior point on this high-rent locus can be replicated by mixing the full-effort endpoint $\tilde{A}_E$ with no-effort actions. The only candidate new extreme action is therefore the full-effort endpoint itself. Economically, flexible effort does not create a new kind of reward; it merely limits how generous the easy full-effort reward can be.

Full effort. At full effort the agent cannot overshoot the target. The relevant contract is therefore again an upper-tail bonus, exactly as in the baseline model. What changes is that full effort must now beat *every* lower effort level, not just zero effort. For an upper-tail bonus this creates two transparent tests: a global test comparing full effort with zero effort and a marginal test asking whether the last unit of effort is worth its cost. Remarkably, for the Gaussian threshold contract these two tests are jointly sufficient.

Lemma OA.3. *Let $w = M\mathbf{1}\{y \geq h - z\sqrt{h}\}$. Then w implements $\eta = h$ in the sense of $(\text{IC}_h^{\text{flex}})$ if and only if*

$$\Delta(h, z) \geq \frac{ch}{M} \quad \text{and} \quad \phi(z) \geq \kappa.$$

The two inequalities measure different ways in which the same standard can fail. The first asks whether the standard generates enough incentive *on average* to make full effort preferable to giving up altogether. The second asks whether it generates enough incentive *at the margin* to stop the agent from shaving the last bit of effort. Phase-level effort checks only the first. Flexible effort checks both.

Lemma OA.4. *Write $\Lambda(h, z) = \sqrt{h} \phi(z)$, let $z^\dagger(h)$ be the unique bar at which $\Delta(h, z) = \Lambda(h, z)$ — it lies in $(0, \sqrt{h}/2)$ — and set*

$$\overline{\mathcal{M}}(h) = \max_z \min\{\Delta(h, z), \Lambda(h, z)\} = \sqrt{h} \phi(z^\dagger(h)). \quad (\text{OA.9})$$

- (i) *Full effort is implementable under flexible effort if and only if $\overline{\mathcal{M}}(h) \geq ch/M$; and $\overline{\mathcal{M}}(h) < \overline{\Delta}(h)$ for every $h > 0$, so flexible effort always tightens feasibility.*
- (ii) *When that inequality is strict, the full-effort slice of $\mathcal{A}^{\text{flex}}(h)$ is the segment of the line $s = (1 - c)h$ joining A_D to $\tilde{A}_E$. The difficult end is unchanged: it is borne by the bar $z_D(h)$ of Section 3.1, at which the global condition binds and the local one is slack, so*

$$r_D(h) = M \Phi(z_D(h) - \sqrt{h}), \quad \pi'(h) = \frac{M \phi(z_D(h))}{\sqrt{h}} > c.$$

The easy end changes: *it is borne by the strictly harder bar $\zeta(h) < z_E(h)$, at which the local condition binds and the global one is slack, so*

$$\tilde{r}_E(h) = M \Phi(\zeta(h)) - ch < r_E(h), \quad \pi'(h) = c. \quad (\text{OA.10})$$

Which of the two conditions of Lemma OA.3 binds is what makes the correction asymmetric, for the reason given in Section 6.1: the difficult standard is already sensitive to effort at the margin, whereas the easy one is not. The magnitudes are worth recording. Feasibility itself changes only modestly: $\overline{\Delta}(h)/\overline{\mathcal{M}}(h) = 1 + h/72 + O(h^2)$. The larger effect is on the rent that the easy action can promise. For example, at $h = 1/4$ and $c = 0.3$, $\tilde{r}_E/r_E$ ranges from about 0.84 near the feasibility floor to 0.98 when the cap is four times that floor; at $h = 1$ the ratio can fall to about 0.68 near the tight end. Flexible effort therefore matters most exactly where continuation rents are scarce and valuable as incentive capital.

No effort. Nothing changes at the no-effort end. A constant payment $w \equiv p$ gives the same expected payment at every effort level, so exerting positive effort only adds cost. Hence every $p \in [0, M]$ implements zero effort and delivers $(r, s) = (p, 0)$. The entire old no-effort segment remains feasible, including $A_B = (0, 0)$ and $A_G = (M, 0)$.

OA.1.3 Extreme points

The review problem cares not about all feasible contracts but about the convex hull of feasible rent-surplus pairs, which is what Figure 4 draws. This section supplies the characterization behind that figure: the relaxed pointwise-effort set lies inside $\text{co}\{A_B, A_D, \tilde{A}_E, A_G\}$, with $\tilde{A}_E$ as in (OA.8) and A_B, A_D, A_G as in Section 3.3, and each of those four vertices is implementable. The one non-obvious step is the low-rent boundary: partial effort never delivers effort more cheaply in rent-per-unit terms than full effort.

The residual question, in rent terms. Could the principal ever prefer to buy *some* continuation effort rather than the full amount, because partial effort requires less rent per unit? Let $\rho(\eta)$ denote the rent under the cheapest relaxed contract implementing η :

$$\rho(\eta) = \hat{P}_{\min}(\eta) - c\eta, \quad \text{with} \quad \rho(h) = M\Phi(z_D(h)) - ch = r_D(h). \quad (\text{OA.11})$$

The relevant chord inequality is equivalent to

$$\frac{\rho(\eta)}{\eta} \geq \frac{\rho(h)}{h} = \frac{r_D(h)}{h} \quad \text{for every } \eta \in (0, h). \quad (\text{OA.12})$$

So the issue has a direct economic interpretation: is full effort the cheapest way to buy effort in rent terms? The answer is yes. Rent per unit falls as the principal buys more effort, so a partial-effort continuation is never an exposed action.

The normalization used in the proof makes this a one-dimensional statement. Writing $\theta = \eta/\sqrt{h}$ and $\kappa = c\sqrt{h}/M$, the minimum rent takes the form

$$\rho(\eta) = M \mathcal{R}_\kappa(\eta/\sqrt{h}), \quad \mathcal{R}_\kappa(\theta) = \Phi(b + \theta) - \Phi(a + \theta), \quad (\text{OA.13})$$

where (a, b) solve the band conditions. The next lemma establishes the desired economies of scale in incentive rent.

Lemma OA.5. *Suppose $\overline{\mathcal{M}}(h) > ch/M$. Then*

$$\widehat{P}_{\min}(\eta) \geq \frac{\eta}{h} \widehat{P}_{\min}(h) \quad \text{for every } \eta \in (0, h);$$

equivalently, by (OA.11), the agent's rent per unit of effort implemented, $\rho(\eta)/\eta$, is smallest at full effort.

Proposition OA.1. *Suppose $\overline{\mathcal{M}}(h) > ch/M$. Then*

$$\text{co } \mathcal{A}^{\text{flex}}(h) = \text{co}\{A_B, A_D, \tilde{A}_E, A_G\},$$

and all four vertices belong to $\mathcal{A}^{\text{flex}}(h)$. In particular, for every $q \in \mathbb{R}$ a maximizer of $s + qr$ over $\mathcal{A}^{\text{flex}}(h)$ can be chosen from these four actions, and no interior effort level is ever used.

This is the payoff from the preceding machinery. Although flexible effort creates a continuum of possible effort levels and a two-dimensional feasible set, none of the interior effort levels is ever chosen by the principal. The economically relevant set still has four vertices. Three are exactly the baseline actions, and the fourth is the same easy continuation with a slightly harder standard.

The proof deliberately uses the relaxed set only as an outer bound. We therefore do not need a complete characterization of incentive compatibility at every interior target.

OA.1.4 Consequences for the analysis

Proposition OA.1 is what licenses the substitution $r_E(h) \mapsto \tilde{r}_E(h)$ announced in Section 6.1. Only the upper breakpoint of (6) moves, to $\tilde{q}_2(h) = (1-c)h/(M - \tilde{r}_E(h)) < q_2(h)$, so the good-stop region opens at a lower shadow value than in the baseline, while the cutoff formula (B.1) is unchanged and the ordered four-region menu of Proposition 2 survives verbatim. Feasibility of full continuation effort is now the strictly stronger requirement $\overline{\mathcal{M}}(h) \geq ch/M$ of Lemma OA.4(i), in place of $\overline{\Delta}(h) \geq ch/M$.

OA.1.5 Flexible effort before the review

So far we have allowed the agent to coast only after the review. The same concern can arise before the review: an agent who is supposed to work until τ may instead work for only part of that interval. The effect is again local and economically transparent. Because the desired first-phase effort is the upper boundary $\eta_1 = \tau$, the principal does not need to make the agent indifferent at the margin; she only needs the marginal return to the last bit of effort to be at least its cost.

Only one new constraint. Let $\eta_1 \in [0, \tau]$ denote total first-phase effort. The review state is $X_\tau \sim N(\eta_1, \tau)$, and the menu promises continuation utility $r(x)$ after review state x . Implementing

full first-phase effort requires $\mathbb{E}_\tau[r(X_\tau)] - c\tau \geq \mathbb{E}_{\eta'_1}[r(X_\tau)] - c\eta'_1$ for every $\eta'_1 \in [0, \tau]$. The baseline model checks only the most distant deviation, $\eta'_1 = 0$. Flexible effort adds the possibility of shaving a small amount of first-phase effort. At the full-effort corner this requires

$$\left. \frac{d}{d\eta_1} \mathbb{E}_{\eta_1}[r(X_\tau)] \right|_{\eta_1=\tau} = \frac{1}{\tau} \text{Cov}_\tau(r(X_\tau), X_\tau) \geq c.$$

Thus the review menu must provide enough continuation-utility sensitivity not only in total, but also at the margin.

The two conditions in the menu's cutoffs. Because the continuation problem still has four extreme actions, the promised-rent schedule is an increasing step function,

$$r(x) = \sum_{i=1}^3 \delta_i \mathbf{1}\{x \geq k_i\}, \quad k_1 < k_2 < k_3,$$

with jumps $\delta_1 = r_D(h)$, $\delta_2 = \tilde{r}_E(h) - r_D(h)$, and $\delta_3 = M - \tilde{r}_E(h)$. Let $z_i = (\tau - k_i)/\sqrt{\tau}$. Then the global and marginal first-phase incentive requirements are

$$\underbrace{\sum_i \delta_i \Delta(\tau, z_i)}_{\text{global}} \geq c\tau, \quad \underbrace{\sum_i \delta_i \Lambda(\tau, z_i)}_{\text{local}} \geq c\tau.$$

The first asks whether the menu makes full effort preferable to zero effort. The second asks whether the last increment of first-phase effort is worth taking. This is exactly the same average-versus-marginal distinction that hardened the easy continuation standard.

The ordered menu survives. Attach multipliers $\lambda \geq 0$ and $\mu \geq 0$ to the global and local first-phase constraints. The review-state index becomes

$$Q(x) = \lambda \left[1 - \exp\left(\frac{\tau}{2} - x\right) \right] + \mu \frac{x - \tau}{\tau} - 1, \quad Q'(x) = \lambda \exp\left(\frac{\tau}{2} - x\right) + \frac{\mu}{\tau} > 0. \quad (\text{OA.14})$$

The important fact is not the extra term itself but its monotonicity. Better review performance still raises the attractiveness of higher-rent actions. Hence the statewise problem still selects among the same four actions in the same order. Flexible first-phase effort changes the cutoffs and introduces a second multiplier, but it does not destroy the ordered menu.

One prediction does change. There is one notable implication. In the baseline model the review-state index is bounded above, so the good-stop region appears only for some parameters. If the marginal first-phase constraint binds, then $\mu > 0$ and the linear term in (OA.14) makes $Q(x)$ unbounded above. The good-stop region must then be nonempty.

Economically, marginal first-phase incentives are especially well supplied by very high review

outcomes. A low termination threshold can create a large difference between working and not working while doing little to reward the *last* increment of effort. A top prize does the opposite: its covariance with performance is large in the upper tail. Freeing first-phase effort therefore shifts incentive provision toward the good-stop prize and makes the full four-region ladder generic whenever the local constraint actually bites.

When nothing changes. The extra local constraint need not bind. Since $\Lambda(\tau, z) \geq \Delta(\tau, z)$ exactly when $z \leq z^\dagger(\tau)$, the global condition automatically implies the local one whenever all cutoffs satisfy $k_i \geq k^\dagger(\tau) = \tau - z^\dagger(\tau)\sqrt{\tau} \in (\tau/2, \tau)$. Because the cutoffs are ordered, it is enough to check the lowest one. Thus if the termination threshold is already sufficiently demanding—roughly above half of expected first-phase output—allowing the agent to coast before the review changes nothing.

A step that does not carry over. There is one limitation. For a single continuation threshold, the global and local conditions are jointly sufficient because the agent’s loss from deviating is unimodal. Before the review the rent schedule can have three jumps, so the marginal-return function is a Gaussian mixture:

$$\frac{d}{d\eta_1} \mathbb{E}_{\eta_1} [r(X_\tau)] = \frac{1}{\sqrt{\tau}} \sum_i \delta_i \phi\left(\frac{k_i - \eta_1}{\sqrt{\tau}}\right).$$

With several peaks, an intermediate effort level can be attractive even when both the global and endpoint-marginal tests pass.

This qualification appears to matter only for long first phases. Numerically the mixture is single-peaked on $[0, \tau]$ for every cutoff configuration we tried at $\tau \leq 4$; a second peak appears by $\tau = 9$ and a third by $\tau = 64$. Thus for short first phases the two conditions above are exact in the cases examined, whereas late reviews require direct checks of intermediate deviations.

A margin left open. Finally, this section asks whether the principal can still implement full first-phase effort. It does not ask whether she would optimally choose less. Allowing the principal to target an interior first-phase effort would introduce a first-order equality before the review as well and would put the band-contract logic on both sides of the review. We leave that additional margin for future work.

OA.1.6 Proofs

OA.1.6.1 Proof of Lemma OA.1

Proof. Consider minimizing $\pi(0) = \int w f_0$ over $0 \leq w \leq M$ subject to the two equality constraints $\int w f_\eta = \pi_w(\eta)$ and $\int w \partial_\eta f_\eta = c$. The feasible set contains w and is the intersection of the weak-* compact ball with two weak-* closed hyperplanes, hence weak-* compact; the objective is weak-* continuous, so a minimizer $\tilde{w}$ exists. Attaching multipliers λ_1, λ_2 to the two constraints gives a Lagrangian linear in w , which is therefore minimized pointwise over $[0, M]$, with $\tilde{w} = M$ exactly

where the index is negative. Dividing that index by $f_\eta(y) > 0$ and using

$$\frac{f_0(y)}{f_\eta(y)} = \exp\left(\frac{\eta^2}{2h} - \frac{\eta y}{h}\right), \quad \frac{\partial_\eta f_\eta(y)}{f_\eta(y)} = \frac{y - \eta}{h},$$

it becomes

$$\Psi(y) = \exp\left(\frac{\eta^2}{2h} - \frac{\eta y}{h}\right) - \lambda_1 - \frac{\lambda_2}{h}(y - \eta),$$

a strictly convex function of y less an affine one, hence strictly convex. Its negativity set is an interval, which is (OA.4), and since a strictly convex function vanishes at most twice the boundary is null, so no randomization is needed. The first two displayed properties hold because $\tilde{w}$ is feasible for the auxiliary problem and the third because it is optimal. Finally, if w satisfies (G $_\eta$) then $\pi_{\tilde{w}}(\eta) - c\eta = \pi_w(\eta) - c\eta \geq \pi_w(0) \geq \pi_{\tilde{w}}(0)$, so $\tilde{w}$ satisfies it too. $\square$

OA.1.6.2 Proof of Lemma OA.2

Proof. That the achievable payments form an interval with the stated endpoints is Lemma OA.1 together with linearity of $\pi(\eta)$ in w on a convex set; the global condition binds at the lower end because lowering $\pi(\eta)$ tightens it. For the upper end let $w^* = M\mathbf{1}\{y \geq \eta - \zeta\sqrt{h}\}$ with $\zeta = \zeta(h)$ and let the agent choose η' , writing $v = (\eta' - \eta)/\sqrt{h}$. Then $\pi(\eta') = M\Phi(\zeta + v)$ and, using $c\sqrt{h} = M\kappa$, his payoff is $M\Phi(\zeta + v) - c\eta - M\kappa v$, whose derivative in v is $M[\phi(\zeta + v) - \kappa]$. Since $\phi(\zeta) = \kappa$ and ϕ is symmetric and unimodal, this is positive exactly on $(-2\zeta, 0)$ and negative outside $[-2\zeta, 0]$. The payoff is therefore maximized at $v = 0$ or at the left endpoint $v = -\theta$, so (IC $^{\text{flex}}_h$) reduces to the single comparison with $\eta' = 0$, namely $D(\theta) := \Phi(\zeta) - \Phi(\zeta - \theta) - \kappa\theta \geq 0$. Now $D(0) = 0$ and $D'(\theta) = \phi(\zeta - \theta) - \kappa$ is positive for $\theta \in (0, 2\zeta)$ and negative for $\theta > 2\zeta$, so $\{\theta \geq 0 : D(\theta) \geq 0\}$ is an interval $[0, \bar{\theta}]$. The hypothesis is exactly $D(\sqrt{h}) \geq 0$, so $\bar{\theta} \geq \sqrt{h}$ and w^* is incentive compatible at every $\eta \leq h$. $\square$

OA.1.6.3 Proof of Lemma OA.3

Proof. Write a deviation as $\eta' = h - u\sqrt{h}$ with $u \in [0, \sqrt{h}]$. The agent's loss from it is $g(u) = M[\Phi(z) - \Phi(z - u)] - cu\sqrt{h}$, so (IC $^{\text{flex}}_h$) says $g \geq 0$ on $[0, \sqrt{h}]$. Now $g(0) = 0$ and $g'(u) = M\phi(z - u) - c\sqrt{h}$. As u rises, $z - u$ falls, so $\phi(z - u)$ rises while $z - u > 0$ and falls thereafter: g' is single-peaked. Given $g'(0) \geq 0$, which is $\phi(z) \geq \kappa$, it follows that g' changes sign at most once and only from positive to negative, so g is unimodal and attains its minimum over $[0, \sqrt{h}]$ at an endpoint. Since $g(0) = 0$, $g \geq 0$ if and only if $g(\sqrt{h}) \geq 0$, which is $M\Delta(h, z) \geq ch$. Conversely, if $\phi(z) < \kappa$ then $g'(0) < 0$ and $g < 0$ just to the right of zero. $\square$

OA.1.6.4 Proof of Lemma OA.4

Proof. By Lemma OA.3 a bar is admissible if and only if $\min\{\Delta(h, z), \Lambda(h, z)\} \geq ch/M$, where $\Lambda(h, z) = \sqrt{h}\phi(z)$, so an admissible bar exists exactly when $\overline{\mathcal{M}}(h) \geq ch/M$.

Where the maximum sits. Substituting $s \mapsto z - s$, $\Delta/\Lambda = h^{-1/2} \int_0^{\sqrt{h}} e^{zs-s^2/2} ds$, whose derivative in z is $h^{-1/2} \int_0^{\sqrt{h}} s e^{zs-s^2/2} ds > 0$; the limits are 0 and ∞ , so the two cross once, at some $z^\dagger(h)$. At $z = 0$ the ratio is $h^{-1/2} \int_0^{\sqrt{h}} e^{-s^2/2} ds < 1$, so $z^\dagger > 0$. At $z = \sqrt{h}/2$, writing $a = \sqrt{h}/2$, the ratio exceeds one exactly when $2\Phi(a) - 1 > 2a\phi(a)$; the difference vanishes at $a = 0$ and has derivative $2a^2\phi(a) > 0$, so $z^\dagger < \sqrt{h}/2$. Hence $\min\{\Delta, \Lambda\}$ equals Δ , increasing, to the left of $z^\dagger$ and Λ , decreasing on $z > 0$, to its right, and the maximum in (OA.9) is attained at $z^\dagger$. Since $\Delta(h, \cdot)$ is strictly unimodal with peak $\bar{\Delta}(h)$ at $\sqrt{h}/2 \neq z^\dagger$, the inequality $\bar{\mathcal{M}}(h) < \bar{\Delta}(h)$ is strict at every $h > 0$.

The two ends. The admissible bars are $z \in [z_D(h), \zeta(h)]$. For the easy end, $z_E(h) > \sqrt{h}/2 > z^\dagger(h)$ gives $\Lambda(h, z_E) < \Delta(h, z_E) = ch/M = \Lambda(h, \zeta(h))$, and Λ is decreasing on $z > 0$, so $\zeta(h) < z_E(h)$. For the difficult end, $\Delta(h, z_D) = ch/M \leq \Delta(h, z^\dagger)$ with $\Delta(h, \cdot)$ increasing to the left of $\sqrt{h}/2$ gives $z_D(h) \leq z^\dagger(h)$, whence $\Lambda(h, z_D) \geq \Delta(h, z_D) = ch/M$: the local condition is slack there and $\pi'(h) = M\phi(z_D)/\sqrt{h} > c$. The rent at an admissible bar is $M\Phi(z) - ch$, strictly increasing in z , so the two ends of the segment are the two ends of the interval of bars, which is (OA.10). At $z_D(h)$ the global constraint binds and the rent collapses to $M\Phi(z_D - \sqrt{h})$ by (1); at $\zeta(h)$ it does not, which is why the easy rent must be computed as $\mathbb{E}_h[w] - ch$ rather than as $\mathbb{E}_0[w]$. $\square$

OA.1.6.5 Proof of Lemma OA.5

Proof. Step 1: the jump at the full-effort face. We first show $\hat{P}_{\min}(h) \leq \liminf_{\eta \uparrow h} \hat{P}_{\min}(\eta)$, equivalently $r_D(h) \leq M\mathcal{R}_\kappa(\sqrt{h})$. The local condition is an equality below the face and an inequality on it, so the constraint set at $\eta = h$ is the larger one. Formally, take $\eta_n \uparrow h$ and let $w_n \in \hat{\mathcal{F}}_h(\eta_n)$ attain $\hat{P}_{\min}(\eta_n)$. The ball $\{0 \leq w \leq M\}$ is weak-* compact, so a subnet converges to some w^* . The maps $\eta \mapsto f_\eta$ and $\eta \mapsto \partial_\eta f_\eta$ are continuous in L^1 , so the constraints pass to the limit: $\pi_{w^*}(h) - ch \geq \pi_{w^*}(0)$ and $\pi'_{w^*}(h) = c$, whence $\pi'_{w^*}(h) \geq c$ and $w^* \in \hat{\mathcal{F}}_h(h)$. Since $\pi_{w^*}(h) = \lim_n \hat{P}_{\min}(\eta_n)$, the claim follows. Equivalently and more directly: any solution of the pair (OA.5) and (OA.7) at $\theta = \sqrt{h}$ satisfies (G_η) at $\eta = h$ by construction and (L_η) with equality, hence is admissible at full effort, so its payment bounds $\hat{P}_{\min}(h)$ from above.

Step 2: a closed form for $\mathcal{R}'_\kappa$. Parametrize the bands satisfying (OA.5) by a , with $b = b(a)$ solving $\phi(b) = \phi(a) - \kappa$, and write $E_x = \phi(x + \theta)/\phi(x) = e^{-\theta x - \theta^2/2}$. Differentiating $\phi(b) = \phi(a) - \kappa$ and using $\phi'(x) = -x\phi(x)$ gives $b'(a) = a\phi(a)/(b\phi(b))$. Write $P(a) = \Phi(b) - \Phi(a)$, $\mathcal{R}(a, \theta) = \Phi(b + \theta) - \Phi(a + \theta)$ and $G(a, \theta) = P(a) - \mathcal{R}(a, \theta) - \kappa\theta$ for the slack in (G_η) . Then

$$P'(a) = \phi(b)b'(a) - \phi(a) = \phi(a) \frac{a - b}{b} < 0,$$

so minimizing the payment means taking the largest admissible a ; call it $a^*(\theta)$, at which (G_η) binds. Substituting $\phi(x + \theta) = \phi(x)E_x$,

$$\mathcal{R}_a = \phi(a) \left(\frac{a}{b} E_b - E_a \right), \quad \mathcal{R}_\theta = \phi(b)E_b - \phi(a)E_a, \quad G_a = \frac{\phi(a)}{b} \left[a(1 - E_b) - b(1 - E_a) \right].$$

Differentiating $G(a^*(\theta), \theta) = 0$ and $\mathcal{R}_\kappa(\theta) = \mathcal{R}(a^*(\theta), \theta)$, and using $G_\theta = -\mathcal{R}_\theta - \kappa$ and $G_a = P' - \mathcal{R}_a$,

$$\mathcal{R}'_\kappa = \mathcal{R}_\theta - \mathcal{R}_a \frac{G_\theta}{G_a} = \frac{P' \mathcal{R}_\theta + \kappa \mathcal{R}_a}{P' - \mathcal{R}_a}.$$

Expanding the numerator and substituting $\kappa = \phi(a) - \phi(b)$, the terms in E_a and E_b collect as $P' \mathcal{R}_\theta + \kappa \mathcal{R}_a = \frac{\phi(a)}{b} (a\phi(a) - b\phi(b))(E_b - E_a)$, so dividing by G_a ,

$$\mathcal{R}'_\kappa(\theta) = \frac{(a\phi(a) - b\phi(b))(E_b - E_a)}{a(1 - E_b) - b(1 - E_a)}. \quad (\text{OA.15})$$

Step 3: signing it. Since $\phi(b) < \phi(a)$ we have $|b| > |a|$, which together with $b > a$ forces $b > 0$; as $\theta > 0$, $E_b < E_a$ and the second factor of the numerator is negative. Differentiating $\pi(\eta') = M[\Phi(b + \theta') - \Phi(a + \theta')]$ twice at $\theta' = 0$ gives $\pi''(\eta) = \frac{M}{h} (a\phi(a) - b\phi(b))$, so the first factor is nonpositive precisely by the second-order condition (S_η). Finally, the denominator equals $bG_a/\phi(a)$ with $\phi(a) > 0$ and $b > 0$; since a^* is the largest a at which $G \geq 0$, and $G(a^*, \theta) = 0$, we have $G_a(a^*, \theta) \leq 0$. The numerator of (OA.15) is therefore nonnegative and the denominator nonpositive, so $\mathcal{R}'_\kappa \leq 0$.

Conclusion. Since $\mathcal{R}_\kappa \geq 0$ and $\mathcal{R}'_\kappa \leq 0$,

$$\frac{d}{d\theta} \left[\frac{\mathcal{R}_\kappa(\theta)}{\theta} \right] = \frac{\theta \mathcal{R}'_\kappa(\theta) - \mathcal{R}_\kappa(\theta)}{\theta^2} \leq 0.$$

Hence for $\theta = \eta/\sqrt{h} < \sqrt{h}$, using (OA.13) and then Step 1,

$$\frac{\rho(\eta)}{\eta} = \frac{M \mathcal{R}_\kappa(\theta)}{\theta \sqrt{h}} \geq \frac{M \mathcal{R}_\kappa(\sqrt{h})}{h} \geq \frac{r_D(h)}{h} = \frac{\rho(h)}{h},$$

which is (OA.12). □

Remark OA.1. *The two signs invoked in Step 3 of the proof are economic rather than technical. That $a\phi(a) \leq b\phi(b)$ is the agent's second-order condition: at the difficult end of the band family the contract would make the target a local minimum of the agent's payoff, and such contracts implement nothing. That $a(1 - E_b) \leq b(1 - E_a)$ says the payment-minimizing band sits on the boundary of the region where the agent still prefers the target to abandoning effort; pushing the bar any further would break that constraint. The two never conflict: at the extreme $a = \zeta(h)$ the band degenerates to an upper-tail bonus, where $a\phi(a) - b\phi(b) = \zeta(h)\kappa > 0$ and (S_η) fails outright, so the binding restriction on a is always (G_η). Numerically, over the entire admissible range $\kappa \in (0, \phi(0))$ and θ up to 12, the global constraint binds at every optimum and (S_η) is strictly slack.*

OA.1.6.6 Proof of Proposition OA.1

Proof. Let $\hat{\mathcal{A}}(h)$ be defined as in (OA.3) with $[\hat{P}_{\min}, \hat{P}_{\max}]$ in place of $[P_{\min}, P_{\max}]$. By (OA.1), $\mathcal{A}^{\text{flex}}(h) \subseteq \hat{\mathcal{A}}(h)$.

We first show $\widehat{\mathcal{A}}(h)$ lies in the quadrilateral with vertices $A_B, A_D, \widetilde{A}_E, A_G$. Because $(\eta, P) \mapsto (r, s)$ is linear and invertible, it is enough to show the feasible (η, P) region lies in the quadrilateral with vertices $(0, 0), (h, \widehat{P}_{\min}(h)), (h, \widehat{P}_{\max}(h))$ and $(0, M)$. The region lies between the graphs of $\widehat{P}_{\min}$ and $\widehat{P}_{\max}$, so it suffices to bound each.

Upper edge. By (OA.6), $\widehat{P}_{\max}(\eta) \leq M\Phi(\zeta(h)) = \widehat{P}_{\max}(h)$, while the upper edge of the quadrilateral is the chord from $(0, M)$ to $(h, \widehat{P}_{\max}(h))$, which lies weakly above that constant because $M \geq \widehat{P}_{\max}(h)$. This step is unconditional.

Lower edge. The lower edge is the chord from $(0, 0)$ to $(h, \widehat{P}_{\min}(h))$, so the requirement is $\widehat{P}_{\min}(\eta) \geq (\eta/h)\widehat{P}_{\min}(h)$, which is Lemma OA.5. By Lemma OA.4, $\widehat{P}_{\min}(h) = M\Phi(z_D(h))$ and $\widehat{P}_{\max}(h) = M\Phi(\zeta(h))$, so the quadrilateral is exactly $\text{co}\{A_B, A_D, \widetilde{A}_E, A_G\}$.

It remains to place the four vertices in $\mathcal{A}^{\text{flex}}(h)$. A_B and A_G are the constant contracts $w \equiv 0$ and $w \equiv M$, covered by the no-effort regime of Section OA.1.2; A_D and $\widetilde{A}_E$ are covered by Lemma OA.4(ii). Hence

$$\text{co}\{A_B, A_D, \widetilde{A}_E, A_G\} \subseteq \text{co } \mathcal{A}^{\text{flex}}(h) \subseteq \text{co } \widehat{\mathcal{A}}(h) = \text{co}\{A_B, A_D, \widetilde{A}_E, A_G\},$$

and the chain closes. □

Remark OA.2. *Because (OA.13) eliminates the horizon, the whole argument can be checked exhaustively rather than sampled. Over the entire admissible range $\kappa \in (0, \phi(0))$ — feasibility forces $\kappa < \phi(0) = 1/\sqrt{2\pi}$, since $\overline{\mathcal{M}}(h) \leq \sqrt{h}\phi(0)$ by (OA.9) — and θ up to 12, which covers every (h, c, M) satisfying $\overline{\mathcal{M}}(h) > ch/M$ with $h \leq 144$: the closed form (OA.15) agrees with numerical differentiation to six decimals, both sign conditions of Step 3 hold at every point, and the jump of Step 1 is strictly positive throughout.*

OA.2 Beyond Gaussian noise

Sections 3–5 use Brownian output. This appendix separates the parts of the analysis that depend only on the geometry of a binary experiment from those that use Gaussian short-horizon asymptotics. Effort remains binary in both phases; pointwise effort is studied in Section 6.1.

Section 6.2 states the separation established here: the phase problem and the ordered fixed-review menu follow from binary-experiment geometry, whereas the review-date results require additional short-horizon properties of the likelihood ratio. Proofs are collected in Section OA.2.6.

OA.2.1 A general phase experiment

Fix a phase of length $h > 0$ and replace the Brownian increment by a primitive binary experiment. The phase produces a public signal valued in a standard Borel space $\mathcal{Y}_h$, distributed P_1^h if the agent works and P_0^h if he shirks. Let $\nu_h = P_0^h + P_1^h$, write $p_j^h = dP_j^h/d\nu_h$, and define the extended likelihood ratios $\Lambda_h = p_1^h/p_0^h$ and $\ell_h = p_0^h/p_1^h$, with the usual conventions at zero densities. Under two-sided evidence these ratios are finite and reciprocal almost surely. Working costs the agent ch

and raises the principal's expected gross output by $m(h)$ relative to shirking; as in Section 2 the shirking baseline is normalized to zero, so $m(h)$ is throughout an *incremental* expected output.⁸

A contract is a measurable $w : \mathcal{Y}_h \rightarrow [0, M]$; let $\mathcal{F}_h$ be the set satisfying $\mathbb{E}_1[w] - \mathbb{E}_0[w] \geq ch$. Writing $\varphi = w/M$ and $\alpha(\varphi) = \mathbb{E}_0[\varphi]$ and $\beta(\varphi) = \mathbb{E}_1[\varphi]$, the contract is a randomized test of shirking against work, of size α and power β . Expected compensation under work is $M\beta$, the incentive spread is $M(\beta - \alpha)$, and the agent's rent when (IC_h) binds is $M\alpha$: *the rent is the cap times the size of the test*. Everything in the phase problem therefore depends on the experiment only through its *receiver operating characteristic*, the compact convex set $\mathcal{R}_h = \{(\alpha(\varphi), \beta(\varphi)) : 0 \leq \varphi \leq 1\} \subset [0, 1]^2$. Let $\beta_h(\cdot)$ be the upper boundary of $\mathcal{R}_h$ — the power of the most powerful test of each size — and set $g_h(\alpha) = \beta_h(\alpha) - \alpha$ and $\delta(h) = ch/M$. Under mutual absolute continuity, three standard facts are all we use. The map β_h is concave and nondecreasing with $\beta_h(0) = 0$ and $\beta_h(1) = 1$; its slope at α is the likelihood ratio at the Neyman–Pearson threshold of that size, so that g_h is concave with $g_h(0) = g_h(1) = 0$ and slope $\Lambda_h - 1$; and $\max_\alpha g_h(\alpha) = \text{TV}_h := \sup_A [P_1^h(A) - P_0^h(A)]$, the total variation distance between the two phase laws.⁹ Thus the paper's implementability threshold becomes $M_I(h, c) = ch/\text{TV}_h$.

The translation into the notation of Section 3 is exact. With $P_j^h = N(jh, h)$ and $m(h) = h$ we have $\Lambda_h(y) = \exp(y - h/2)$ and, parameterizing sizes by $\alpha = \Phi(z - \sqrt{h})$,

$$\beta_h(\alpha) = \Phi(z), \quad g_h(\alpha) = \Delta(h, z), \quad \bar{\Delta}(h) = \text{TV}_h = 2\Phi\left(\frac{\sqrt{h}}{2}\right) - 1.$$

The incentive index $\Delta(h, \cdot)$ of Section 3 is therefore the ROC gap function, Figure 1 is a picture of a concave function vanishing at both ends, and its two roots $z_D(h)$ and $z_E(h)$ are the two ends of a superlevel set of that function. This is why the figure has the shape it does, and it is the reason the next lemma requires no parametric distributional form.

Lemma OA.6. *Fix $h > 0$ and $M > 0$, and suppose P_1^h and P_0^h are mutually absolutely continuous.*

- (i) $\mathcal{F}_h \neq \emptyset$ if and only if $ch \leq M \text{TV}_h$.
- (ii) If $ch < M \text{TV}_h$, every contract that minimizes expected compensation over $\mathcal{F}_h$ is a likelihood-ratio test $w = M\mathbf{1}\{\Lambda_h > t\}$, with fractional payment permitted on $\{\Lambda_h = t\}$, whose size is $\alpha_L(h) := \inf\{\alpha : g_h(\alpha) \geq \delta(h)\}$.
- (iii) Writing $\alpha_H(h) := \sup\{\alpha : g_h(\alpha) \geq \delta(h)\}$, so that $[\alpha_L(h), \alpha_H(h)]$ is the compact superlevel interval, and

$$r_L(h) = M\alpha_L(h), \quad r_H(h) = M\alpha_H(h), \quad \Sigma(h) = m(h) - ch,$$

⁸If shirking itself produces output $m_0(h) > 0$, that baseline cancels from the within-phase work–shirk comparison but not automatically from the comparison with stopping. The formulas below then require the stop payoff to be adjusted unless the same baseline accrues under the outside option. In the Brownian model $m_0 \equiv 0$ and $m(h) = h$.

⁹Where Λ_h has atoms, β_h is piecewise linear and “slope” means the appropriate one-sided derivative; the maximizer of g_h is then a vertex and need not be a size at which Λ_h equals one. We therefore define roots by infima and suprema below. Smooth comparative statics require the additional differentiability conditions stated where they are used.

the set of continuation rents $r = \mathbb{E}_1[w] - ch$ attainable while implementing work is exactly $[r_L(h), r_H(h)]$, and the full-effort frontier is the horizontal segment $s = \Sigma(h)$ over that interval, along which the principal's payoff $u = \Sigma(h) - r$ falls one for one with the rent conceded.

Whenever work is implementable, let $\beta_L(h) = \beta_h(\alpha_L(h))$ and use the paper's no-review notation for the cheapest work-inducing contract: $R^{\text{NR}}(h; c, M) = r_L(h)$ and $V^{\text{NR}}(h; c, M) = \Sigma(h) - r_L(h) = m(h) - M\beta_L(h)$. The contract is profitable relative to stopping exactly when $V^{\text{NR}}(h; c, M) \geq 0$.

Lemma OA.6 is the implementability half of Lemma 2, Lemma 1, and Lemma 3 at once, with $\bar{\Delta}(h)$ read as TV_h . What the Gaussian added was the extra conclusion, in Lemma 1, that a likelihood-ratio test is a *cutoff in the phase signal*. That conclusion requires monotonicity of Λ_h but is not needed for the continuation-rent geometry. The distinct monotonicity condition imposed below concerns Λ_τ and orders review actions in the observed review score X_τ .

OA.2.2 Five assumptions behind the incentive ladder

Fix a review date $\tau \in (0, T)$ and write $h = T - \tau$. The first phase generates the experiment (P_1^τ, P_0^τ) and the continuation phase the experiment (P_1^h, P_0^h) . We maintain throughout that the output process has stationary and independent increments, so that the continuation problem depends on the review state only through the remaining horizon h ; this is what licenses the reduction at the head of Section 3.3, and it is a property of the Brownian model that has nothing to do with normality.¹⁰

Consider the following conditions.

Assumption 5 (Two-sided evidence). *The two laws are mutually absolutely continuous in both phases: $P_1^t \sim P_0^t$ for $t \in \{\tau, h\}$.*

Assumption 6 (Slack detection). *$ch < M \text{TV}_h$ and $c\tau < M \text{TV}_\tau$.*

Assumption 7 (Profitable continuation). *$\Sigma(h) = m(h) - ch > 0$.*

Assumption 8 (Monotone review evidence). *Λ_τ is nondecreasing in the performance measure X_τ .*

Assumption 9 (Randomization). *Either the law of Λ_τ under P_1^τ is atomless, or the principal may publicly randomize between adjacent review actions.*

Assumption 5 says that no realization of the phase signal is conclusive: neither work nor shirking can be ruled out by any observation. Assumption 6 is the strict phase-by-phase analogue of the implementability condition (2). Unlike Assumption 4, it concerns implementability only; profitability is imposed separately below. Assumption 7 holds automatically in the Brownian model, where $\Sigma(h) = (1 - c)h > 0$ because $c < 1$; with a general output technology it is a restriction. Assumption 8 is the monotone likelihood-ratio property, which the Gaussian delivers through (4). Assumption 9 is what makes the review cutoffs well defined and the primal problem attain the dual bound; it is automatic for the Brownian review signal.

¹⁰Without it the phase laws depend on X_τ , so r_L , r_H and Σ become functions of the review state. The statewise reduction of Proposition 2 survives, but the ordering does not: the ladder ranks four actions by the rent they deliver, and that ranking is informative only if the rents themselves do not move with the state. Models with learning about project quality, or with effort-dependent volatility, must confront this separately.

Proposition OA.2. *Suppose Assumptions 5–9 hold at (τ, h) . Then*

$$0 < r_L(h) < r_H(h) < M, \quad (\text{OA.16})$$

the convex hull of the review action set $\mathcal{A}(h) = \mathcal{A}_0(h) \cup \mathcal{A}_1(h)$ has exactly the four extreme points

$$A_B = (0, 0), \quad A_L = (r_L(h), \Sigma(h)), \quad A_H = (r_H(h), \Sigma(h)), \quad A_G = (M, 0),$$

and the phase reduction, pre-review Lagrangian, and fixed-review action ordering hold after the substitutions $\bar{\Delta} \mapsto \text{TV}$ and $(1-c)h \mapsto \Sigma(h)$. The review-date derivatives and endpoint results require the additional assumptions studied below. In particular the support function is

$$\mathcal{V}(q, h) = \max\{0, \Sigma(h) + qr_L(h), \Sigma(h) + qr_H(h), qM\}, \quad (\text{OA.17})$$

the index is $Q_\lambda = \lambda(1 - \ell_\tau) - 1$, the optimal menu selects A_B, A_L, A_H, A_G in that order as Q_λ crosses the three breakpoints

$$q_0 = -\frac{\Sigma(h)}{r_L(h)} < 0 < q_2 = \frac{\Sigma(h)}{M - r_H(h)}, \quad (\text{OA.18})$$

and the fixed-review value is

$$V^R(\tau) = \inf_{\lambda \geq 0} \left\{ m(\tau) - \lambda c\tau + \mathbb{E}_1[\mathcal{V}(Q_\lambda, h)] \right\}. \quad (\text{OA.19})$$

For a positive supporting multiplier away from score ties, the menu has at most four ordered regions in X_τ ; public randomization implements the corresponding mixture at an atom or a zero-multiplier tie.

The whole of the ladder is thus one geometric fact: bounded compensation applied to a signal that is informative but never conclusive generates a nondegenerate *interval* of continuation rents compatible with continuation work, strictly inside the interval $[0, M]$ of rents compatible with stopping. The two stops are the two extreme rents, the two work actions are the two ends of the work interval, and a monotone review score ranks all four by the rent they deliver.

OA.2.3 Three identities that replace the Gaussian formulas

Proposition OA.2 preserves the structure of the analysis; it does not preserve the closed forms, which used the normal distribution function throughout. Three exact identities replace them, and each specializes to a display already derived in Sections 3–5.

Rent is cost divided by evidence. Recall that $\beta_L(h)$ is the normalized expected payment under work at the cheapest contract, and define the *evidence ratio* of the reward event,

$$\bar{\Lambda}_h = \frac{\beta_L(h)}{\alpha_L(h)} = \frac{\mathbb{E}_1[\varphi_L]}{\mathbb{E}_0[\varphi_L]} > 1. \quad (\text{OA.20})$$

If the generalized test is implemented with an independent uniform draw, the last ratio is the work-to-shirk likelihood ratio averaged over the event on which the cap is paid. This interpretation also covers a fractional payment at a likelihood-ratio atom. Since $M[\beta_L - \alpha_L] = ch$ and $R^{\text{NR}}(h) = r_L(h) = M\alpha_L(h)$,

$$r_L(h) = \frac{ch}{\bar{\Lambda}_h - 1}, \quad \text{and dually} \quad M - ch - r_H(h) = \frac{ch}{\bar{\Lambda}_h - 1}, \quad (\text{OA.21})$$

where $\beta_H(h) = \beta_h(\alpha_H(h))$ and $\bar{\Lambda}_h^- = (1 - \alpha_H)/(1 - \beta_H)$ is the shirk-to-work odds ratio of the event on which payment is withheld at the rent-maximizing contract. Equation (OA.21) isolates the role of limited liability: the rent conceded per unit of effort cost is the reciprocal of the excess evidence carried by the rewarded event. The two ratios in (OA.21) are *average* likelihood ratios over the reward and withholding events, not the marginal likelihood ratio at the cutoff, which is the distinct object ℓ_h appearing below. Letting $M \rightarrow \infty$ drives $\alpha_L(h)$ to zero and $\bar{\Lambda}_h$ up to $\beta'_h(0^+) = \Lambda_h^{\max}$, the essential supremum of the likelihood ratio, so $\lim_{M \rightarrow \infty} R^{\text{NR}}(h; c, M) = ch/(\Lambda_h^{\max} - 1)$. The Gaussian case is $\Lambda_h^{\max} = \infty$, in which the limit is zero; this is the step in the proof of Lemma 2 that drives $r_D(h) \rightarrow 0$ as $M \rightarrow \infty$. When the available likelihood ratio is bounded — for example under some coarse monitoring technologies — agency rent does *not* vanish as the cap loosens. The closing discussion of Section 3 should accordingly be read as a statement about signals with unbounded likelihood ratios: it is the boundedness of the *evidence*, and only derivatively the boundedness of the *prize*, that keeps rents alive. A positive rent creates scope for a review to recycle continuation utility as incentives, but does not by itself prove that a review dominates the no-review contract.

The symmetry identity. The relation $r_L(h) + r_H(h) = M - ch$, which in the Brownian model follows from $z_E(h) = \sqrt{h} - z_D(h)$, is a symmetry property of the experiment.

Proposition OA.3. *Under Assumption 5 the following are equivalent.*

- (i) $r_L(h) + r_H(h) = M - ch$ for every (c, M) with $0 < ch < M \text{TV}_h$.
- (ii) $\mathcal{R}_h$ is invariant under the anti-diagonal reflection $(\alpha, \beta) \mapsto (1 - \beta, 1 - \alpha)$.
- (iii) The experiment obtained by interchanging the two hypotheses, (P_0^h, P_1^h) , is Blackwell equivalent to (P_1^h, P_0^h) .

A sufficient condition is the existence of a measurable map σ with $\sigma_{\#}P_1^h = P_0^h$ and $\sigma_{\#}P_0^h = P_1^h$; in particular the identity holds for every location family with a symmetric, everywhere positive noise density, with $\sigma(y) = m - y$ at mean shift m . It fails, for instance, when effort shifts a Poisson arrival rate.

The no-review derivative and the tight-cap limit. Let $\ell_h = 1/\beta'_h(\alpha_L(h))$ be the likelihood ratio at the binding bar, keep $\lambda_h^{\text{NR}} = 1/(1 - \ell_h)$ as in Section 5.1, and write $\dot{\beta}_h = M \partial_h \beta_h(\alpha)|_{\alpha=\alpha_L(h)}$ for the value of lengthening the phase at a fixed size.

Proposition OA.4. *Suppose m is differentiable at h and $\beta_h(\alpha)$ is differentiable in (h, α) at $\alpha_L(h)$, with $\beta'_h(\alpha_L(h)) > 1$. Then*

$$\frac{dR^{\text{NR}}(h)}{dh} = \frac{c - \dot{\beta}_h}{\beta'_h(\alpha_L(h)) - 1}, \quad \frac{dV^{\text{NR}}(h)}{dh} = m'(h) - \lambda_h^{\text{NR}}[c - \ell_h \dot{\beta}_h]. \quad (\text{OA.22})$$

In the Brownian model $m' \equiv 1$ and $\ell_h \dot{\beta}_h = M\phi(z_D(h) - \sqrt{h})/(2\sqrt{h})$, so (OA.22) is (C.3).

Now let the cap fall to the feasibility boundary. The two roots merge at the maximizer α_h^* of g_h , and the marginal likelihood ratio there approaches one precisely when g_h has no kink at its peak. Writing $\varepsilon(h) = d \log \text{TV}_h / d \log h$ for the elasticity of detection with respect to the phase length, the envelope theorem applied to $\text{TV}_h = \max_\alpha g_h(\alpha)$ gives $\partial_h \beta_h(\alpha_h^*) = \text{TV}'_h$ and hence the following.

Corollary OA.1. *Suppose the differentiability conditions of Proposition OA.4 hold in a neighborhood of the feasibility boundary, TV_h is differentiable at h , and $\text{ess inf}\{\Lambda_h : \Lambda_h > 1\} = 1$, so that the likelihood ratio takes values arbitrarily close to one from above. Then $\ell_h \uparrow 1$ and $\lambda_h^{\text{NR}} \rightarrow \infty$ as $M \downarrow M_I(h, c)$, and*

$$\varepsilon(h) < 1 \implies \lim_{M \downarrow M_I(h, c)} \frac{dV^{\text{NR}}(h)}{dh} = -\infty, \quad \varepsilon(h) > 1 \implies \text{the limit is } +\infty.$$

Moreover $\varepsilon(h) < 1$ if and only if TV_h/h is strictly decreasing at h .

The condition $\varepsilon(h) < 1$ is the property established for the Gaussian in the proof of Lemma 2, where $\bar{\Delta}(h)/h$ is shown to be strictly decreasing. The inequality $2\Phi(a) - 1 > a\phi(a)$ in (C.6), proved in Appendix C, is its Gaussian instance: with $a = \sqrt{h}/2$, $\varepsilon(h) = a\phi(a)/(2\Phi(a) - 1)$. Thus the same elasticity condition underlies both calculations. For a general experiment family it must be imposed explicitly. The regularity condition is also substantive. A two-point signal with likelihood ratios $\{1.75, 0.5\}$ has a gap around one and its ROC has a vertex at the total-variation maximizer; the marginal likelihood ratio on the lower branch is stuck at 1.75, and λ_h^{NR} remains bounded however tight the cap. For such signals the tight-cap mechanism discussed in Section 5.1 — the two incentive-compatible thresholds nearly coincide and the shadow price of incentives explodes — simply does not operate. What drives that mechanism is the *continuity* of the likelihood ratio around one, not its Gaussian form.

OA.2.4 Timing: diffusive versus hard evidence

The results on the review *date* are far less robust than the results on the review *menu*. Throughout this subsection the output process has stationary independent increments, $m(h) = m_1 h$, and the experiment family is continuous in total variation at zero: $\text{TV}_h \rightarrow 0$ as $h \downarrow 0$. Stochastic continuity

alone is not enough for the last property, so it is imposed explicitly. As in the baseline analysis, we also maintain that the full-horizon no-review contract is strictly implementable and profitable, $V^{\text{NR}}(T; c, M) > 0$.

Lemma OA.7. $\text{TV}_{h_1+h_2} \leq \text{TV}_{h_1} + \text{TV}_{h_2}$ for all $h_1, h_2 > 0$, and

$$\Theta_0 := \lim_{h \downarrow 0} \frac{\text{TV}_h}{h} = \sup_{h > 0} \frac{\text{TV}_h}{h} \in [0, +\infty],$$

the limit existing.

The second short-horizon index is the limit of the evidence ratio (OA.20). Fix a cap and cost for which all sufficiently short phases are implementable, and suppose $\Lambda_0 := \lim_{h \downarrow 0} \bar{\Lambda}_h$ exists in $[1, \infty]$, so that by (OA.21) the minimum rent per unit of time converges, $r_L(h)/h \rightarrow c/(\Lambda_0 - 1)$, with the convention $c/0 = +\infty$. Say that the family is *diffusive at zero* if $\Lambda_0 = 1$ and that it *carries hard evidence at zero* if $\Lambda_0 > 1$: in the first case the cheapest reward event capable of supplying the required incentive spread has an average likelihood ratio converging to one, while in the second its evidential content remains bounded away from one. Combining, the no-review value per unit of time has the limit

$$\nu := \lim_{h \downarrow 0} \frac{V^{\text{NR}}(h)}{h} = m_1 - c - \frac{c}{\Lambda_0 - 1} \in [-\infty, \infty).$$

The number ν governs whether a vanishing continuation phase is valuable at the terminal endpoint. The initial endpoint instead depends on the two tails of the short-horizon likelihood ratio and on the utility spread available at the review.

Proposition OA.5. *If $M\Theta_0 < c$, then work is implementable over no phase of positive length, so no contract with any number of reviews implements effort anywhere. If $M\Theta_0 > c$, work is implementable over every sufficiently short phase, and if in addition TV_h/h is strictly decreasing and continuous there is a unique maximal implementable phase length $\bar{h}(M)$, the root of $M\text{TV}_h = ch$.*

Proposition OA.5 is where the Brownian model is special. With Brownian noise $\text{TV}_h \sim \sqrt{h/2\pi}$ and $\Theta_0 = \infty$, so short phases are always implementable, whatever the cap: effort cost accumulates linearly in the phase length while distinguishability grows like $\sqrt{h}$, and shortening the detection window therefore buys detection power. For a pure finite-intensity arrival experiment, by contrast, $\Theta_0 < \infty$ and feasibility becomes the *scale-free* restriction $M\Theta_0 > c$, which no subdivision of the horizon can relax. The feasibility index Θ_0 and the evidence ratio Λ_0 capture distinct properties; a technology combining diffusive and jump signals can have $\Theta_0 = \infty$ while also carrying hard rare-event evidence.

The endpoint results require separate conditions, stated next.

Proposition OA.6. *Suppose work is strictly implementable in a neighborhood of T , $V^{\text{NR}}(T; c, M) > 0$, and $r_L(t)$ and $V^{\text{NR}}(t; c, M)$ are locally Lipschitz at $t = T$.*

(i) (Early review.) Suppose there are $\kappa > 0$ and $C < \infty$ and events E_τ in the review σ -field with

$$P_1^\tau(E_\tau) - P_0^\tau(E_\tau) \geq \kappa\tau, \quad P_1^\tau(E_\tau) \leq C\tau$$

for all small τ , and suppose $\kappa[M - r_L(T)] > c$. Then

$$\liminf_{\tau \downarrow 0} \frac{V^R(\tau) - V^{\text{NR}}(T; c, M)}{\tau} > -\infty,$$

so (C.10) fails.

(ii) (Late review.) Suppose $\nu > 0$. Let B_τ be the reward event in a cheapest work-inducing contract for a phase of length τ , on the probability space augmented by any public randomization used at a likelihood-ratio atom. Suppose there are events $E_\tau \subseteq B_\tau^c \cap \{\Lambda_\tau > 1\}$ such that $P_1^\tau(E_\tau) > P_0^\tau(E_\tau)$ and $P_1^\tau(E_\tau) \rightarrow p > 0$ as $\tau \uparrow T$. Then

$$V^R(T - h) \geq V^{\text{NR}}(T - h; c, M) + \nu p h + o(h),$$

so the first-order equivalence $V^R(T - h) = V^{\text{NR}}(T - h; c, M) + o(h)$ of Lemma C.1 fails and, using Lemma 4,

$$\limsup_{h \downarrow 0} \frac{V^{\text{NR}}(T; c, M) - V^R(T - h)}{h} \leq \limsup_{h \downarrow 0} \frac{V^{\text{NR}}(T; c, M) - V^{\text{NR}}(T - h; c, M)}{h} - \nu p.$$

In particular, if the two left derivatives exist, then $V_-^{R'}(T) \leq \partial_T V^{\text{NR}}(T; c, M) - \nu p < \partial_T V^{\text{NR}}(T; c, M)$.

These are separate sufficient conditions. Hard upper-tail evidence helps satisfy the early-review event condition and can make ν finite, but neither the displayed incentive-capital inequality nor $\nu > 0$ follows from hard evidence alone.

The mechanism in part (ii) runs in the opposite direction from the Gaussian case. The Gaussian proof of (C.3) shows that a vanishing continuation phase contributes nothing at first order, and it does so by observing that $r_L(h)$ vastly exceeds the continuation surplus $\Sigma(h)$ when h is small: continuation work is ruinously rent-expensive over a short phase, so the two work actions are dominated on all but a vanishing range of the review index. Under hard evidence the ranking reverses, $r_L(h)$ and $\Sigma(h)$ are both of order h , and if $\nu > 0$ the low-rent continuation is strictly better than the bad stop over a fixed range of review states. A short continuation phase then has *first-order* value, which lowers the upper left slope at T — and the left derivative when it exists — and therefore makes an interior review easier, not harder, to justify at that endpoint.

The mechanism in part (i) is different. The Gaussian argument writes obedience as $\mathbb{E}_1[(r - r_L)(1 - \ell_\tau)] \geq c\tau$ and uses that $|1 - \ell_\tau|$ is of order $\sqrt{\tau \log(1/\tau)}$ outside an event of probability $o(\tau)$, so that buying $c\tau$ of incentive requires utility dispersion of order $\sqrt{\tau}/\sqrt{\log(1/\tau)}$, which is large

relative to τ . The general form of that step is: for every $\theta \in (0, 1)$,

$$\mathbb{E}_1|r - r_L| \geq \frac{c\tau - MG_\tau(\theta)}{\theta}, \quad G_\tau(\theta) := \mathbb{E}_1[(|1 - \ell_\tau| - \theta)^+].$$

Under the maintained strict profitability of the low-rent continuation at T , the principal's loss controls this utility dispersion exactly as in the proof of Proposition 4(i). Hence (C.10) holds whenever one can choose $\theta_\tau \downarrow 0$ with $MG_\tau(\theta_\tau) \leq c\tau/2$. In the Brownian model $\theta_\tau = 2\sqrt{\tau \log(1/\tau)}$ does this and returns the rate in the appendix. The absolute value in G_τ is essential: incentives may be supplied by punishing highly informative evidence of shirking as well as by rewarding evidence of work, and both channels must be shut down for the bound to apply. In a finite-intensity arrival model in which short phases are feasible, $G_\tau(\theta)$ is of order τ for every θ below the likelihood-ratio jump carried by a single arrival. No vanishing θ_τ then satisfies $MG_\tau(\theta_\tau) \leq c\tau/2$, so this argument cannot establish the Gaussian endpoint result.

Proposition OA.6(i) is what qualifies the organizational reading of Proposition 4(i) in Section 6.2: it bounds the cost of a short probation by its length when a rare event — a missed milestone, a lost account, a trial readout — retains substantial evidential content as the window shrinks. Note what it does not deliver. The incentive-capital and profitability conditions remain necessary for probation to be *optimal*; ruling out an infinitely steep initial loss is not the same as making an early review worthwhile.

OA.2.5 Which rungs the ladder has: the two tails of the likelihood ratio

Proposition OA.2 says the optimal menu climbs $A_B \rightarrow A_L \rightarrow A_H \rightarrow A_G$. It does not say that all four rungs are reached. At a positive supporting multiplier the Brownian bad stop is always reached, because arbitrarily bad news is available; the good stop is reached only when incentives are expensive enough. That asymmetry is not general. The following likelihood-ratio bounds characterize which outer actions are reached.

The relevant primitives are the two tails of the review likelihood ratio,

$$\ell_\tau^{\max} = \text{ess sup} \frac{dP_0^\tau}{dP_1^\tau}, \quad \Lambda_\tau^{\max} = \text{ess sup} \frac{dP_1^\tau}{dP_0^\tau}, \quad \text{both in } (1, \infty]. \quad (\text{OA.23})$$

The first measures the strongest available evidence of shirking, the second the strongest available evidence of work. Since $Q_\lambda = \lambda(1 - \ell_\tau) - 1$ is an increasing transformation of the score, the closed essential range of the index is $[\lambda(1 - \ell_\tau^{\max}) - 1, \lambda(1 - 1/\Lambda_\tau^{\max}) - 1]$, and whether it reaches the outer breakpoints of (OA.18) is a question about (OA.23).

Proposition OA.7. *Under the hypotheses of Proposition OA.2, with a positive supporting multiplier $\lambda > 0$:*

(i) *the bad stop is strictly optimal on a set of positive probability if and only if*

$$\lambda(\ell_\tau^{\max} - 1) > \frac{\Sigma(h)}{r_L(h)} - 1, \quad (\text{OA.24})$$

which holds for every $\lambda > 0$ when $\ell_\tau^{\max} = \infty$;

(ii) the good stop is strictly optimal on a set of positive probability if and only if

$$\lambda \left(1 - \frac{1}{\Lambda_\tau^{\max}}\right) > 1 + \frac{\Sigma(h)}{M - r_H(h)}, \quad (\text{OA.25})$$

that is, λ must exceed $(1 + \Sigma(h)/(M - r_H(h)))\Lambda_\tau^{\max}/(\Lambda_\tau^{\max} - 1)$. This reduces to $\lambda > 1 + \Sigma(h)/(M - r_H(h))$ when $\Lambda_\tau^{\max} = \infty$ and diverges as $\Lambda_\tau^{\max} \downarrow 1$.

At equality in either condition, the corresponding stop can be used only on a likelihood-ratio atom at which it ties the adjacent continuation action; whether it is used then depends on tie-breaking and the aggregate obedience requirement.

The conditions reflect the opportunity cost of stopping. The good stop pays the maximal prize and abandons production; the principal chooses it only after evidence of work strong enough to justify surrendering the continuation, and (OA.25) says how strong. The bad stop pays nothing and abandons production; that requires evidence of shirking strong enough to justify the same sacrifice, which is (OA.24). A monitoring technology that cannot generate one kind of evidence cannot make the corresponding rung strictly optimal; at a boundary atom the rung may still appear through tie-breaking.

OA.2.6 Proofs for the beyond-Gaussian extension

Proof of Lemma OA.6. Part (i) restates $\sup_\varphi(\beta - \alpha) = \text{TV}_h$. For (ii), minimizing expected compensation is minimizing β subject to $\beta - \alpha \geq \delta(h)$; the constraint binds, since scaling φ down by $\delta(h)/(\beta - \alpha)$ preserves $0 \leq \varphi \leq 1$, restores equality and strictly lowers β , and among binding tests the Neyman–Pearson lemma selects a likelihood-ratio test. Concavity and continuity of g_h make its superlevel set a compact interval; its left endpoint α_L is the smallest feasible size and hence gives the cheapest binding test. For (iii), maximizing β subject to the same constraint also yields a binding test: if the spread is strictly larger than $\delta(h)$, mixing the test with the constant test $\varphi \equiv 1$ until the spread binds strictly raises β . After the substitution $\psi = 1 - \varphi$, this is the mirror-image problem $\mathbb{E}_0\psi - \mathbb{E}_1\psi \geq \delta(h)$; its solution withholds payment where Λ_h is smallest and has size $\alpha_H(h)$. Convex combinations of the two extreme binding tests remain binding and sweep every intermediate size. Therefore the feasible values of β , and hence of $r = M\beta - ch$, form exactly $[r_L, r_H]$. Finally $u + r = m(h) - M\beta + M\beta - ch = \Sigma(h)$ for every work contract, so total surplus is constant at $\Sigma(h)$ across the interval and the principal’s payoff falls one for one with rent. $\square$

Proof of Proposition OA.2. Assumption 6 makes $\{g_h \geq \delta(h)\}$ a nondegenerate interval, so $\alpha_L(h) < \alpha_H(h)$, and it makes $\alpha_H(h) < 1$ because $g_h(1) = 0 < \delta(h)$; hence $r_L < r_H < M$. Assumption 5 gives $\alpha_L(h) > 0$: a test of size zero has $\beta = 0$ under mutual absolute continuity, hence $g_h(0) = 0 < \delta(h)$. This is (OA.16), and with Lemma OA.6(iii) it identifies the four extreme points of the convex hull of the two segments. The statewise Lagrangian is $\Sigma(h)i(x) + Q_\lambda(x)r(x)$, exactly (5), which is affine in

r with slope Q_λ ; the four candidate rents $0 < r_L < r_H < M$ are therefore selected in increasing order of Q_λ , with the breakpoints obtained by equating the corresponding affine functions in (OA.17), and Assumption 7 orders them as in (OA.18). Assumption 8 makes Q_λ nondecreasing in X_τ , so the regions are intervals of performance. Assumption 9, with Assumption 6 supplying a strictly feasible menu, convexifies the aggregate constraint set and delivers attainment of (OA.19). No other property of the two experiments is used. $\square$

Proof of Proposition OA.3. Complementing a test, $\varphi \mapsto 1 - \varphi$, shows that $\mathcal{R}_h$ is always symmetric about $(\frac{1}{2}, \frac{1}{2})$; interchanging the hypotheses maps (α, β) to (β, α) and so reflects $\mathcal{R}_h$ in the diagonal. Composing the two reflections gives the anti-diagonal reflection, whence (ii) $\Leftrightarrow$ (iii), two binary experiments being Blackwell equivalent exactly when they share an ROC region. (ii) $\Rightarrow$ (i): the reflection preserves $\beta - \alpha$ and maps the upper boundary to itself, so it exchanges the two endpoints of the binding chord, giving $\alpha_H = 1 - \beta_L = 1 - \alpha_L - \delta(h)$. (i) $\Rightarrow$ (ii): mutual absolute continuity makes β_h and g_h continuous. For every $\delta \in (0, \text{TV}_h)$, statement (i) gives $\alpha_H(\delta) = 1 - \alpha_L(\delta) - \delta = 1 - \beta_h(\alpha_L(\delta))$. Thus the anti-diagonal reflection of the left endpoint of every superlevel interval is its right endpoint. As δ varies, these endpoints trace the two sides of the upper ROC boundary; continuity covers any plateau at the maximum and the endpoints at $\delta = 0$. Hence the upper boundary, and therefore the ROC region beneath it, is invariant under the reflection. For the sufficient condition, $\mathbb{E}_0[\varphi \circ \sigma] = \mathbb{E}_1[\varphi]$ and $\mathbb{E}_1[\varphi \circ \sigma] = \mathbb{E}_0[\varphi]$, so $1 - \varphi \circ \sigma$ has size $1 - \beta(\varphi)$ and power $1 - \alpha(\varphi)$. $\square$

Proof of Proposition OA.4. $V^{\text{NR}}(h) = m(h) - M\beta_h(\alpha_L(h))$ and $R^{\text{NR}}(h) = M\alpha_L(h)$, with α_L determined by $M[\beta_h(\alpha_L) - \alpha_L] = ch$. Differentiating the constraint gives $\dot{\beta}_h + (\beta'_h - 1)M\alpha'_L = c$, which is the first display; then $\frac{d}{dh}[M\beta_h(\alpha_L)] = \dot{\beta}_h + \beta'_h M\alpha'_L = (\beta'_h c - \dot{\beta}_h)/(\beta'_h - 1) = (c - \ell_h \dot{\beta}_h)/(1 - \ell_h)$. For the Brownian specialization, holding $\alpha = \Phi(z - \sqrt{h})$ fixed means holding $z - \sqrt{h}$ fixed, so $\partial_h \beta_h = \phi(z_D)/(2\sqrt{h})$, while $\ell_h = \phi(z_D - \sqrt{h})/\phi(z_D)$. $\square$

Proof of Corollary OA.1. As $M \downarrow M_I(h, c)$, the lower root converges to a maximizer α_h^* of g_h . The likelihood-ratio condition and the maintained differentiability imply $\beta'_h(\alpha_L(h)) \downarrow 1$, and therefore $\ell_h \uparrow 1$ and $\lambda_h^{\text{NR}} \rightarrow \infty$. The envelope theorem gives $\partial_h \beta_h(\alpha_h^*) = \text{TV}'_h$, so $c - \ell_h \dot{\beta}_h \rightarrow c - M_I(h, c) \text{TV}'_h = c[1 - \varepsilon(h)]$. The two divergence claims now follow from (OA.22); the finite term $m'(h)$ does not affect the limit. Finally, $d \log(\text{TV}_h/h)/d \log h = \varepsilon(h) - 1$, which proves the last equivalence. $\square$

Proof of Lemma OA.7. The time- $(h_1 + h_2)$ increment is a measurable function of the pair of independent phase increments, so the data-processing inequality and the tensorization bound for total variation give the first claim. For the second, fix $h > 0$ and $0 < s < h$, write $h = ns + r$ with $n = \lfloor h/s \rfloor$ and $0 \leq r < s$, and apply subadditivity: $\text{TV}_h \leq n \text{TV}_s + \text{TV}_r$, so $\text{TV}_h/h \leq \text{TV}_s/s + \text{TV}_r/h$. The maintained total-variation continuity makes $\text{TV}_r \rightarrow 0$ as $s \downarrow 0$, so letting $s \downarrow 0$ along an arbitrary sequence gives $\text{TV}_h/h \leq \liminf_{s \downarrow 0} \text{TV}_s/s$. Taking the supremum over h yields $\sup_h \text{TV}_h/h \leq \liminf_{s \downarrow 0} \text{TV}_s/s \leq \limsup_{s \downarrow 0} \text{TV}_s/s \leq \sup_h \text{TV}_h/h$. $\square$

Proof of Proposition OA.5. By Lemma OA.6(i), a phase of length h is implementable exactly when $\text{TV}_h/h \geq c/M$. Lemma OA.7 shows that the left-hand side attains its supremum in the limit $h \downarrow 0$, which proves the first two claims. Under strict monotonicity and continuity it crosses c/M at most once; it crosses exactly once because its limit at zero is larger than c/M and $\text{TV}_h/h \leq 1/h \rightarrow 0$ as $h \rightarrow \infty$. $\square$

Proof of Proposition OA.6. (i) Consider the two-region menu assigning $A_L(h)$, $h = T - \tau$, off E_τ and the good stop A_G on E_τ . Its incentive supply is $[M - r_L(h)] [P_1^\tau(E_\tau) - P_0^\tau(E_\tau)] \geq \kappa [M - r_L(h)] \tau > c\tau$ for all small τ , since $r_L(h) \rightarrow r_L(T)$; pre-review obedience is an inequality, so the menu is feasible as it stands. Relative to assigning $A_L(h)$ everywhere, the loss is at most $[\Sigma(h) - r_L(h) + M] P_1^\tau(E_\tau) \leq \text{const} \cdot C\tau$. Hence $V^R(\tau) \geq m(\tau) + V^{\text{NR}}(T - \tau; c, M) - O(\tau) \geq V^{\text{NR}}(T; c, M) - O(\tau)$.

(ii) Start from the menu that uses only the two stops, assigning A_G on the cheapest-contract reward event B_τ and A_B off it. It delivers the agent the rent profile of the no-review contract at horizon τ , so obedience binds exactly and its value is $V^{\text{NR}}(\tau; c, M)$. Reassign $A_L(h)$ in place of A_B on the event E_τ in the statement. Obedience is relaxed, since the rent rises by $r_L(h)$ on a set with $P_1^\tau(E_\tau) > P_0^\tau(E_\tau)$, and the principal's surplus changes by $[\Sigma(h) - r_L(h)] P_1^\tau(E_\tau) = V^{\text{NR}}(h) P_1^\tau(E_\tau)$. With $V^{\text{NR}}(h) = \nu h + o(h)$ and $\nu > 0$ this is $\nu ph + o(h) > 0$, giving the stated bound. Subtracting it from Lemma 4, dividing by h , and taking upper limits gives the slope inequality in the statement. $\square$

Proof of Proposition OA.7. By (OA.18) the bad stop is strictly optimal at x exactly when $Q_\lambda(x) < q_0$, where $q_0 = -\Sigma(h)/r_L(h)$, that is when $\ell_\tau(x) > (\lambda - 1 - q_0)/\lambda$. Such x have positive probability if and only if $\ell_\tau^{\max} > (\lambda - 1 - q_0)/\lambda$, that is $\lambda(\ell_\tau^{\max} - 1) > -1 - q_0 = \Sigma(h)/r_L(h) - 1$. If $\ell_\tau^{\max} = \infty$ this holds for every $\lambda > 0$. Similarly the good stop is strictly optimal when $Q_\lambda(x) > q_2 = \Sigma(h)/(M - r_H(h))$, that is when $\ell_\tau(x) < 1 - (1 + q_2)/\lambda$. Since the essential infimum of ℓ_τ is $1/\Lambda_\tau^{\max}$, such x have positive probability if and only if $1/\Lambda_\tau^{\max} < 1 - (1 + q_2)/\lambda$, which is (OA.25). Equality in either comparison is a pointwise tie and can matter only when the boundary likelihood ratio has an atom of positive probability. $\square$

OA.3 Random review dates

OA.3.1 The attainable set and duality

For $\tau \in (0, T]$ let $\mathfrak{A}_\tau$ be the measurable menus $a_\tau(x) \in \mathcal{A}(T - \tau)$ with $\mathcal{A}(0) = \{(r, 0) : r \in [0, M]\}$, so that the date- T menu is a bounded payment schedule on the terminal signal, and let

$$\mathcal{B}_\tau = \{(J(\tau, a), P(\tau, a)) : a \in \mathfrak{A}_\tau\} \subset \mathbb{R}^2. \quad (\text{OA.26})$$

At the zero-date endpoint set $\mathcal{B}_0 = \{b : \exists \tau_n \downarrow 0, b_n \in \mathcal{B}_{\tau_n}, b_n \rightarrow b\}$, and put $\mathcal{S} = \bigcup_{\tau \in [0, T]} \mathcal{B}_\tau$ and $\mathcal{C} = \text{co}(\mathcal{S})$. The zero-date convention is innocuous: every point of $\mathcal{B}_0$ has zero incentive surplus and payoff at most $V^{\text{NR}}(T; c, M)$, which the ordinary no-review contract already attains at $\tau = T$.

Lemma OA.8. *For every $\tau \in (0, T]$, $\mathcal{B}_\tau$ is compact and convex, and $\mathcal{S}$ is compact. If a lottery over finitely many menus at a common positive date τ generates (J, P) , some single deterministic menu at that date generates the same pair.*

Proof. Fix τ and write $h = T - \tau$. The action set $\mathcal{A}(h)$ is compact and uniformly bounded, $0 \leq r \leq M$ and $0 \leq s \leq T$. For a measurable selection $a = (r, s)$ the two coordinates of (J, P) are integrals of the bounded measurable vector

$$v_\tau(x, a) = \left(r[g_1^\tau(x) - g_0^\tau(x)], (s - r)g_1^\tau(x) \right) \quad (\text{OA.27})$$

plus the constant $(-c\tau, \tau)$. The correspondence $F_\tau : x \mapsto v_\tau(x, \mathcal{A}(h))$ is measurable, compact-valued and integrably bounded, and $\mathcal{B}_\tau$ is the translate by $(-c\tau, \tau)$ of its Aumann integral: every menu induces a selection, and conversely Filippov's implicit function theorem turns any measurable selection of F_τ into a measurable menu. That integral is compact, and is convex by Lyapunov's theorem because the measure $[g_0^\tau + g_1^\tau]dx$ is atomless — even though $\mathcal{A}(h)$ is a union of two segments.

The same argument purifies. For two menus a^1, a^2 and $p \in [0, 1]$, the atomless vector measure $\mu(B) = \int_B [v_\tau(x, a^1) - v_\tau(x, a^2)]dx$ has compact convex range containing 0 and $\mu(\mathbb{R})$, hence $p\mu(\mathbb{R})$; using a^1 on a set B_p with $\mu(B_p) = p\mu(\mathbb{R})$ and a^2 elsewhere reproduces both integrals. Finite lotteries follow by induction.

Compactness across dates: for $\tau_n \rightarrow \tau > 0$, $g_j^{\tau_n} \rightarrow g_j^\tau$ in L^1 , and $h \mapsto \mathcal{A}(h)$ has compact values and closed graph on $[0, T]$ — under Assumption 4, $r_D(h)$ and $r_E(h)$ are continuous on $(0, T]$ and $r_D(h) \rightarrow 0$, $r_E(h) \rightarrow M$ as $h \downarrow 0$ (Lemma A.1), so $\mathcal{A}(h) \rightarrow \mathcal{A}(0)$ in the Hausdorff metric. Representing a_n by the measure $m_n(x)dx \delta_{a_n(x)}(da)$ with $m_n = (g_0^{\tau_n} + g_1^{\tau_n})/2$, uniform boundedness and tightness give a weakly convergent subsequence whose limit disintegrates as $m(x)dx \kappa_x(da)$ with κ_x supported on $\mathcal{A}(T - \tau)$. The weights $g_j^{\tau_n}/m_n$ are bounded by two and converge locally uniformly, so both moments in (OA.27) converge. The barycenter is a selection of $\text{co } F_\tau$, whose Aumann integral coincides with that of F_τ for atomless measures, and the Filippov step converts it into a deterministic menu. Hence $b \in \mathcal{B}_\tau$.

If instead $\tau_n \downarrow 0$, the total variation between the two review-score laws is $2\Phi(\sqrt{\tau_n}/2) - 1 \rightarrow 0$, so rents bounded by M give $J \rightarrow 0$; and every action satisfies $s - r \leq (1 - c)h_n - r_D(h_n)$ for large n , so $P \leq \tau_n + (1 - c)h_n - r_D(h_n) \rightarrow V^{\text{NR}}(T; c, M)$. Every point of $\mathcal{B}_0$ is thus $(0, p)$ with $p \leq V^{\text{NR}}(T; c, M)$, and $(0, V^{\text{NR}}(T; c, M)) \in \mathcal{B}_0$. Compactness of $\mathcal{S}$ follows. $\square$

Since a lottery takes probability-weighted averages of deterministic pairs, and any finite convex combination of elements of $\mathcal{S}$ is implemented by a lottery over dates (merging duplicate dates by the lemma, and replacing a $\mathcal{B}_0$ point by $(0, V^{\text{NR}}(T; c, M)) \in \mathcal{B}_T$), Problem (SR) is exactly

$$V^S = \max\{p : (j, p) \in \mathcal{C}, j \geq 0\}, \quad \text{and likewise} \quad V^R(\tau) = \max\{p : (j, p) \in \mathcal{B}_\tau, j \geq 0\}. \quad (\text{OA.28})$$

Strict feasibility holds under Assumption 4: since $M > \underline{M}(T, c) \geq M_I(T, c)$ and $M_I(\cdot, c)$ is strictly increasing (Lemma 2), a bounded threshold reward supplies strictly more than $c\tau$ at every $\tau \in (0, T]$.

With $\mathcal{L}^*(\tau, \lambda) \equiv \max_{(j,p) \in \mathcal{B}_\tau} \{p + \lambda j\}$, strong duality for these one-constraint concave programs therefore gives

$$V^D = \sup_{\tau \in [0, T]} \inf_{\lambda \geq 0} \mathcal{L}^*(\tau, \lambda), \quad V^S = \inf_{\lambda \geq 0} \sup_{\tau \in [0, T]} \mathcal{L}^*(\tau, \lambda), \quad (\text{OA.29})$$

the second because a linear functional has the same maximum on $\mathcal{S}$ and on $\text{co}(\mathcal{S})$. A minimizing multiplier exists: strict feasibility supplies $(j_+, p_+) \in \mathcal{C}$ with $j_+ > 0$, so the dual objective diverges as $\lambda \rightarrow \infty$. Substituting (10) shows $\mathcal{L}^*(\tau, \cdot)$ is the pointwise problem of Section 4.2, so at every date and multiplier a maximizing menu can be selected from the four ordered actions of Proposition 2; the only difference is that a lottery uses one common multiplier across all dates in its support.

Proof of Proposition 8. The inequality $V^D \leq V^S$ is the minimax inequality applied to (OA.29), and also follows because a degenerate Γ reproduces every deterministic date.

Let λ^* minimize the dual and let $\mathcal{S}^* = \{(j, p) \in \mathcal{S} : p + \lambda^* j = V^S\}$, which is nonempty and compact. The exposed face $\mathcal{F}^* = \{(j, p) \in \mathcal{C} : p + \lambda^* j = V^S\}$ equals $\text{co}(\mathcal{S}^*)$: by Carathéodory every point of $\mathcal{C}$ is a convex combination of at most three points of $\mathcal{S}$, and such an average attains the maximum of a linear functional only if all weight sits on maximizers.

If $\lambda^* = 0$, strong duality gives $V^S = \max\{p : (j, p) \in \mathcal{C}\}$, so an optimal primal pair lies in $\mathcal{F}^*$ and is feasible; since every point of $\mathcal{S}^*$ has payoff V^S , some element of $\mathcal{S}^*$ has $j \geq 0$, and one date suffices. If $\lambda^* > 0$, any optimal primal pair has $p = V^S$ and $p + \lambda^* j \leq V^S$, hence $j = 0$: aggregate obedience binds and $(0, V^S) \in \mathcal{F}^*$. That face lies on the line $p + \lambda^* j = V^S$, so $(0, V^S)$ is a convex combination of at most two points of $\mathcal{S}^*$ — one if zero is itself an attainable incentive coordinate there, otherwise two whose incentive coordinates bracket zero. Merging a repeated date by Lemma OA.8 gives at most two distinct dates.

Suppose $V^S > V^D$. No optimum can sit at a single date: support at the zero-date endpoint is impossible because every pair there has payoff at most $V^{\text{NR}}(T; c, M) \leq V^D < V^S$, and at a single $\tau \in (0, T]$ the aggregate pair lies in the convex compact $\mathcal{B}_\tau$ and is therefore attained by one feasible deterministic contract, giving $V^D \geq V^S$. Neither supported surplus can vanish, since a supported pair with $J_i = 0$ lies in $\mathcal{S}^*$, hence has $P_i = V^S$, and is again a feasible deterministic contract. As their convex combination is zero the two surpluses have opposite signs; label the two pairs b_+ and b_- so that $J_+ > 0 > J_-$, and solve $pJ_+ + (1 - p)J_- = 0$ for (11).

For the test, necessity follows by substituting (11) into $V^S = p(b_+, b_-)P_+ + [1 - p(b_+, b_-)]P_-$, which is $\Pi(b_+, b_-)$ by (12) and equals $V^S > V^D$, so (13) holds. Conversely, two pairs satisfying (13) can be mixed with probability $p(b_+, b_-)$ to give zero aggregate surplus and payoff above V^D , so $V^S > V^D$; they cannot share a date, since the lemma would then replace the mixture by a feasible deterministic menu at that date. $\square$

Equivalently, each attainable pair generates an affine dual payoff $P + \lambda J$; the surplus branch slopes up and the deficit branch slopes down, and (13) says two such lines cross above V^D . At an interior minimizing multiplier zero must lie in the convex hull of the slopes of the exposed lines: an exposed line of slope zero is a pair that is obedient on its own, and deterministic timing attains V^S ;

under a strict gain there is none, and the zero-slope combination of two opposite-sloped lines is the optimal lottery.

OA.3.2 Two dates under pointwise effort

Throughout this subsection continuation actions after the review are the pointwise-feasible ones of Section 6.1, so $\tilde{A}_E$ replaces A_E .

Proof of Proposition 9. Dynamic reduction. Before τ_1 no review can occur, so the agent receives no information about the realized branch, and because the review-score law and the linear effort cost depend on the pre- τ_1 path only through total effort, any deviation is summarized by $e \in [0, \tau_1]$. If the early review occurs his utility is $U_1(e)$. If it does not, he learns the date is τ_2 ; starting from e , feasibility confines cumulative effort at τ_2 to $[e, e + d]$, so his continuation value is $\hat{U}_2(e)$ in (14). Choosing e is therefore worth $pU_1(e) + (1 - p)\hat{U}_2(e)$. On path $e = \tau_1$, and obedience after non-arrival requires cumulative effort τ_2 , which is $U_2(\tau_2) = \hat{U}_2(\tau_1)$. Given that, full effort before τ_1 is optimal if and only if $pU_1(\tau_1) + (1 - p)U_2(\tau_2) \geq pU_1(e) + (1 - p)\hat{U}_2(e)$ for every e , which rearranges to (15).

The bound. If $U_2(\tau_2) \geq U_2(m)$ on $[0, \tau_2]$ then $\hat{U}_2(e) \leq U_2(\tau_2)$, so $D_2(e) \geq 0$; the same condition applied at $m = \tau_1$ gives $U_2(\tau_2) = \hat{U}_2(\tau_1)$. Where $D_1(e) \geq 0$, (15) holds for every p . Where $D_1(e) < 0$ it is equivalent to $p[D_2(e) - D_1(e)] \leq D_2(e)$, that is $p \leq D_2(e)/[D_2(e) - D_1(e)]$; deterring all deviations is therefore $p \leq \bar{p}$ in (16).

Dominance. Under full effort the lottery is worth $V(p) = P_2 + p(P_1 - P_2)$, affine and increasing in p , so $V(p) > V^{D, \text{pw}}$ if and only if $p > \underline{p}$. A feasible strict improvement exists exactly when $(\underline{p}, 1]$ meets $[0, \bar{p}]$, which is (17). $\square$

Under a shape condition the bound reduces to a comparison of marginal incentive returns. Suppose the early branch locally favors less effort, $U'_1 < 0$ on $[0, \tau_1]$, and that after non-arrival the late-branch optimum is always to work at the maximal rate,

$$\hat{U}_2(e) = U_2(e + d), \quad U'_2(e + d) > 0 \quad \text{on } [0, \tau_1]. \quad (\text{OA.30})$$

Write $\alpha = -U'_1 > 0$ and $\beta(\cdot) = U'_2(\cdot + d) > 0$, so that $-D_1(e) = \int_e^{\tau_1} \alpha$ and $D_2(e) = \int_e^{\tau_1} \beta$, and hence

$$\frac{D_2(e)}{D_2(e) - D_1(e)} = \frac{\int_e^{\tau_1} \beta}{\int_e^{\tau_1} [\alpha + \beta]} = \frac{\int_e^{\tau_1} \rho w}{\int_e^{\tau_1} w}, \quad \rho \equiv \frac{\beta}{\alpha + \beta}, \quad w \equiv \alpha + \beta, \quad (\text{OA.31})$$

a weighted average of ρ on $[e, \tau_1]$.

Corollary OA.2. *Under (OA.30), if ρ is weakly increasing the binding deviation is $e = 0$ and $\bar{p} = D_2(0)/[D_2(0) - D_1(0)]$; if ρ is weakly decreasing the binding deviation is infinitesimal coasting just before τ_1 and $\bar{p} = \rho(\tau_1)$. In the latter case (17) reads*

$$\beta(\tau_1)(P_1 - V^{D, \text{pw}}) > \alpha(\tau_1)(V^{D, \text{pw}} - P_2). \quad (\text{OA.32})$$

Proof. Deleting the lowest part of $[e, \tau_1]$ raises a weighted average of a weakly increasing ρ , so the infimum in (OA.31) is at $e = 0$; if ρ is weakly decreasing the average falls as the left end is deleted, so the infimum is approached as $e \uparrow \tau_1$, and continuity gives $\rho(\tau_1)$. Substituting into $\underline{p} < \bar{p}$ and clearing positive denominators gives (OA.32). $\square$

Randomization is thus profitable when the payoff gained per unit of incentive deficit on the attractive branch exceeds the payoff sacrificed per unit of incentive capacity on the supporting branch.

OA.3.3 Numerical construction

All figures are direct evaluations of the Gaussian formulas; no simulation is used. Set $T = 1.5$, $c = 0.72$, $M = 2.60$ throughout.

Phase-level example. Optimizing first in λ and then in τ gives $\tau^D = 0.861662$, $\lambda^D = 1.824076$ and $V^D = 0.07850435$. Reversing the order gives the minimizing common multiplier $\lambda^* = 1.969842$, at which the maximized Lagrangian has two global maximizers in the review date:

	review date	principal payoff P	incentive surplus J
early branch (b_-)	0.094868	0.12577784	-0.02182519
later branch (b_+)	0.816614	-0.02796598	0.05622360

The later branch carries the incentive surplus, so it is b_+ . Equation (11) places probability $p(b_+, b_-) = 0.279635$ on it and the remaining 0.720365 on the early branch; aggregate obedience binds, and $V^S = 0.08278565$, so $V^S - V^D = 0.00428130 > 0$. At a candidate (τ, λ) these values require only the breakpoints $q_0, 0, q_2$ of (6), their score cutoffs from (B.1) with $k_j = +\infty$ for unreached breakpoints, and the region probabilities

$$\pi_i^j(\tau, \lambda) = \Phi\left(\frac{\bar{k}_i - j\tau}{\sqrt{\tau}}\right) - \Phi\left(\frac{k_i - j\tau}{\sqrt{\tau}}\right), \quad j \in \{0, 1\}, \quad (\text{OA.33})$$

from which $P = \tau + \sum_i (s_i - r_i) \pi_i^1$ and $J = \sum_i r_i [\pi_i^1 - \pi_i^0] - c\tau$.

Pointwise example. The same two dates are used, so $h_1 = 1.405132$ and $h_2 = 0.683386$. With $z_D(h)$ the difficult root of $\Delta(h, z) = ch/M$ and $\zeta(h)$ the positive root of $\phi(\zeta) = c\sqrt{h}/M$ from (8),

$$r_D(h) = M\Phi(z_D(h) - \sqrt{h}), \quad \tilde{r}_E(h) = M\Phi(\zeta(h)) - ch,$$

giving $z_D = 0.03576003$, $\zeta = 0.62451798$, $r_D = 0.32538914$, $\tilde{r}_E = 0.89633126$ on the early branch and $z_D = -0.64321880$, $\zeta = 1.05396840$, $r_D = 0.18406886$, $\tilde{r}_E = 1.72849548$ on the late branch. The menus use cutoffs

	τ_i	k_{i1}	k_{i2}	k_{i3}
early	0.094868	-0.044984	0.878889	0.970519
late	0.816614	0.404840	1.146958	1.503501

so that the rent schedule is the step function with jumps $\delta_{i1} = r_D(h_i)$, $\delta_{i2} = \tilde{r}_E(h_i) - r_D(h_i)$ and $\delta_{i3} = M - \tilde{r}_E(h_i)$, and

$$U_i(e) = \sum_{j=1}^3 \delta_{ij} \Phi\left(\frac{e - k_{ij}}{\sqrt{\tau_i}}\right) - ce, \quad U'_i(e) = \frac{1}{\sqrt{\tau_i}} \sum_{j=1}^3 \delta_{ij} \phi\left(\frac{e - k_{ij}}{\sqrt{\tau_i}}\right) - c. \quad (\text{OA.34})$$

Under full effort the region probabilities are (0.32489, 0.66965, 0.00322, 0.00223) on the early branch and (0.32431, 0.31834, 0.13375, 0.22359) on the late branch, whence $P_1 = 0.1330054377$ and $P_2 = 0.0319887113$.

For the sequential check, $U_1(\tau_1) = 0.15828980$ and $U_1(0) = 0.18420688$, so $D_1(0) = -0.02591708$, and $U'_1 \in [-0.27499, -0.27207]$ on $[0, \tau_1]$: the early branch by itself strictly favors less pre-review effort throughout. On the late branch $U_2(\tau_2) = 0.28317018$ while U'_2 has a unique zero on $[0, \tau_2]$ at $m_0 = 0.35917$, a strict minimum with $U_2(m_0) = 0.21524173$; hence $U_2(\tau_2)$ is the global maximum and the late branch is self-incentive-compatible. Because U_2 has only this interior minimum, $\hat{U}_2(e) = \max\{U_2(e), U_2(e + d)\}$, with the two terms crossing once at $e_c = 0.00215$. At $e = 0$ the agent prefers not to add effort after learning the date, so $\hat{U}_2(0) = U_2(0) = 0.25991515$ and $D_2(0) = 0.02325503$. The deviation-specific bound $q(e) = D_2(e)/[D_2(e) - D_1(e)]$ is increasing on $[0, e_c]$ and on $[e_c, \tau_1]$, so the binding deviation is $e = 0$ and

$$\bar{p} = q(0) = \frac{0.02325503}{0.02325503 + 0.02591708} = 0.47293135.$$

Since the phase-level problem checks fewer deviations both before and after the review, $V^{D, \text{pw}} \leq V^D = 0.07850435$ and $\underline{p} \leq 0.46047462 < \bar{p}$. At $p = 0.465$ the worst deviation still loses $0.465D_1(0) + 0.535D_2(0) = 0.00039 > 0$, and the contract is worth $0.465P_1 + 0.535P_2 = 0.07896149$, exceeding V^D by 0.00045714.