

Contracting and Disclosure of Managerial Actions*

Kohei Kawamura[†]

August 2026

Abstract

This paper studies joint contract and disclosure design when disclosure concerns a productive managerial action chosen by the principal. Disclosure changes the agent's effort response, which changes the principal's action choice and hence the information the disclosure conveys. Every optimal disclosure policy is one-sided: the action for which effort has the larger effect on expected output always generates the signal that induces effort, while the other does so with weakly lower probability. Strict partial disclosure can be optimal. Full disclosure can be suboptimal because revealing the action with the smaller output return to effort discourages effort even when the effort would increase total surplus. If the principal can commit to her action rule, full disclosure is weakly optimal.

Keywords: moral hazard, incentive contracts, disclosure, managerial actions, information design.

JEL codes: D23, D82, D83, D86.

1 Introduction

When management is choosing among strategic actions that affect employees' incentives, it must decide both how to reward employees and what to tell them about the action it will take. Warner Bros. Discovery faced this problem during its 2025 strategic review. Its board of directors was considering whether to proceed with the planned separation into Warner Bros. and Discovery Global, sell the company as a whole, pursue separate transactions for either business, or combine a Warner Bros. transaction with a spin-off of Discovery Global. During this process, the company established a \$38.7 million transaction bonus pool for selected employees.¹ Employees were therefore being

*I thank Shinichi Hirota, Akifumi Ishihara, James Malcomson, Margaret Meyer, David Myatt, Marco Ottaviani and seminar participants at Oxford and Waseda for helpful comments and discussions. This work was supported by JSPS KAKENHI Grant Numbers JP19H01471 and 25K00623. All errors are my own.

[†]Graduate School of Economics, Faculty of Political Science and Economics, Waseda University. kkawamura@waseda.jp

¹WBD announced in October 2025 that its board was considering the planned separation of Warner Bros. and Discovery Global, a sale of the whole company, separate transactions involving either business, and an alternative structure combining a Warner Bros. transaction with a Discovery Global spin-off. On 3 December, before the board selected a transaction, the Compensation Committee adopted a \$38.7 million Transaction Bonus Program for selected key employees. The board ultimately approved a transaction under which Warner Bros. would be acquired by Netflix following the separation of Discovery Global. See WBD's 2025 third-quarter filings and December 2025 Form 8-K.

given incentives to contribute before management had determined which strategic course of action it would ultimately pursue, even though that choice would itself affect employees' roles, organizational affiliation, and future prospects. Similar combinations of restricted transaction information and employee incentives appear in other recent deals.² The relevant disclosure problem is not only what employees should learn about the prospects of a transaction. It is what they should learn about the action management itself will take while that action remains under management's control.

The same issue arises in outsourcing, relocation, restructuring, and technological change. Employees' return to effort can depend on what management chooses to do. Disclosure of a managerial action can therefore change employees' anticipated effort response. Since effort need not affect all managerial actions in the same way, that response can change the relative payoff from the available actions and hence management's choice. Since that choice is itself the object of disclosure, the change in behaviour alters what employees can infer from the disclosure policy.³ This raises the question studied in this paper: how should a principal design incentives and disclosure when disclosure concerns her own action and that action remains subject to her choice?

This paper studies this question in an agency model. A principal chooses between two productive actions, A and B . Binary effort raises the probability of high verifiable output under either action, but its marginal effect on output is larger under A than under B . Only action B gives the principal an additional private benefit. The principal privately observes the size of this benefit before choosing her action, and effort can increase the benefit. Thus effort affects the principal differently across actions. It raises output more under A , while under B it also raises the private benefit. Before observing the size of the private benefit, the principal commits to an output-contingent incentive contract and a disclosure policy that generates internal signals about the chosen action. She then observes the private benefit and chooses A or B . The agent observes the resulting signal before choosing whether to exert effort. The signal cannot enter the contract.

We first demonstrate that an arbitrary disclosure policy can be reduced to two internal signals associated with the principal's action choice. After one signal the agent exerts effort, while after the other he does not. Every optimal policy inducing a nondegenerate equilibrium has a one-sided form. Action A always generates the effort-inducing signal, while B generates it with weakly lower probability. A strictly partial policy therefore reveals action B with positive probability and pools B with A otherwise. The numerical example establishes that such strict partial disclosure can be uniquely optimal. By contrast, if the principal can commit to how her action depends on the private benefit and jointly optimize the contract and disclosure policy, full disclosure is weakly optimal.

Increasing the probability that B generates the signal followed by effort can induce productive

²Confluent's sale process restricted non-public due diligence information to parties under confidentiality agreements, while IBM negotiated executive retention and transition arrangements before the merger agreement of 7 December 2025. Spirit AeroSystems' merger agreement with Boeing allowed a cash retention bonus programme of up to \$50 million. See the companies' respective proxy statements. These filings establish the coexistence of restricted transaction information and employee incentives, not the causal mechanism studied in this paper.

³Management research documents both links separately. Communication and employee interpretation affect responses to mergers and strategic change (Schweiger and DeNisi, 1991; Sonenshein and Dholakia, 2012; Angwin et al., 2016). Anticipated employee mobility and retention also affect acquisition decisions (Younge et al., 2015; Chen et al., 2021). Section 2 discusses these papers in relation to the model.

effort after B more often. It also makes B more attractive to the principal and lowers the cutoff below which she chooses A , so B is chosen for a wider range of Z . The lower cutoff makes the principal choose B for some values of Z for which A would create more total surplus. This offsets part of the gain from inducing effort under B . The cutoff also changes the relative frequencies of A and B behind the signal followed by effort, and hence what that signal means for the return to effort. This changes the bonus required to induce effort. The bonus in turn affects the cutoff because a larger bonus raises the principal’s expected bonus payment more under A , where effort has a larger effect on output. The cutoff and bonus must therefore satisfy a fixed point. Partial disclosure balances the gain from inducing effort with positive probability when B is chosen against the surplus loss from the resulting reduction in the cutoff.

Section 2 relates the paper to the literature. Section 3 presents the model. Section 4 reduces arbitrary disclosure policies to two internal signals and establishes the one-sided form. Section 5 studies joint contract and disclosure design and gives a partial-disclosure example. Section 6 examines commitment to the managerial action and contractible signals. Section 7 presents extensions and robustness. Section 8 concludes. Appendix A contains the proofs.

2 Related literature

Two agency papers are especially close for different reasons. Blanes i Vidal and Möller (2007) show that sharing hard information about project prospects can affect employee motivation and thereby the leader’s project choice, and they also study the joint choice of information sharing and incentives. Their disclosed information, however, is generated independently of the leader’s project choice. Here the managerial action itself is disclosed. Disclosure therefore changes the agent’s effort, which changes the principal’s action choice and hence what the signal conveys. This feedback generates the fixed-point problem at the centre of our analysis.

Tan (2023) shows that partial transparency about a principal’s privately chosen ability to monitor the agent’s effort can improve incentives. We study disclosure of a productive managerial action whose payoff depends on both the principal’s private information and the agent’s effort. Moreover, the contract and disclosure policy are jointly designed: the bonus affects which action the principal chooses and therefore what the disclosure signal conveys.

Management research documents separately the two behavioural links that the present model joins. On one side, communication changes how employees respond to mergers and strategic change: Schweiger and DeNisi (1991) study employee communication following a merger in a longitudinal field experiment. On the other side, anticipated employee retention changes which transactions firms undertake: using a natural experiment, Younge et al. (2015) find that constraints on employee mobility increase acquisition likelihood.⁴ These papers establish the two directions separately. The

⁴Related evidence includes Sonenshein and Dholakia (2012) on employee interpretation and engagement in strategic change implementation, Angwin et al. (2016) on communication approaches and merger and acquisition outcomes, and Chen et al. (2021) on the Inevitable Disclosure Doctrine, acquisition likelihood, and post-acquisition retention of key employees.

present model joins them when the strategic choice itself is disclosed: communication changes worker behaviour, anticipated worker behaviour changes the principal’s strategic action, and that action changes what subsequent communication conveys.

A broad literature studies disclosure in moral hazard environments with fixed payoff-relevant states. Jehiel (2015) gives general reasons why full transparency may be suboptimal in organizations. Catonini and Stepanov (2026) provide sufficient conditions under which full disclosure is optimal and apply them to familiar principal-agent specifications. Kwak (2022) studies optimal information disclosure when a state changes the output technology. Curello et al. (2026) distinguish the efficiency gain from tailoring effort to a state from the agency cost saving obtained by aggregating the agent’s incentive constraints across states. These papers take the payoff-relevant state as fixed. Here the principal chooses the action being disclosed, so changing disclosure changes the behaviour that generates the agent’s information. When the principal can instead commit to how her action depends on private information and jointly optimize the contract, full disclosure is weakly optimal.

The paper also relates to the joint design of information and incentive systems. Lewis and Sappington (1997) study information management when an agent acquires planning information before supplying hidden effort. Li and Yang (2020) characterize optimal incentive contracts when the monitoring technology is itself designed. Sobel (1993) studies information control in a principal-agent relationship. The monitoring technology in these papers provides information to the principal or concerns information used in contracting. In this paper, the internal signal is sent to the agent and concerns a productive action chosen by the principal.

Several papers study information design when strategic behaviour changes the distribution that information is about. Boleslavsky and Kim (2021) consider Bayesian persuasion when an agent’s hidden effort changes the distribution of an underlying state. Babichenko et al. (2026) characterize information structures through which a principal observes an agent’s strategic action. Their direction of information is the reverse of ours. In this paper, the signal informs the agent about the principal’s action, and the principal’s response changes the distribution of that action. The joint design problem therefore requires the agent’s incentive condition and the principal’s best response to be jointly consistent.

The model also relates to work in which the designer is also a player in the game. Baliga et al. (1997) and Baliga and Sjöström (1999) study implementation when the planner or principal participates in the game rather than simply committing to an outcome rule. Jost (1996) studies an informed principal who takes a further action after contracting and compares commitment and noncommitment to that action. In this paper, the principal can commit to the contract and disclosure policy but not to the productive action she chooses after acquiring private information. The relevant restriction is therefore on commitment to the principal’s subsequent action, rather than on her ability to commit to the contract and disclosure policy.

Finally, the paper builds on the standard moral hazard, coordination, and persuasion frameworks of Holmström (1979), Myerson (1982), and Kamenica and Gentzkow (2011). Here the principal both designs the disclosure policy and chooses the action it describes, so Bayes’ rule and the principal’s

best response must be jointly consistent.

3 Model

There is a risk-neutral principal and a risk-neutral agent. The principal chooses an action

$$a \in \{A, B\},$$

and the agent subsequently chooses effort

$$e \in \{0, 1\}.$$

Exerting effort costs $c > 0$ to the agent. Output is also binary, $y \in \{0, 1\}$,⁵ and

$$\Pr(y = 1 \mid a, e) = p^0 + e\Delta_a.$$

We assume

$$0 < \Delta_B < c < \Delta_A, \quad 0 \leq p^0 \leq 1 - \Delta_A.$$

Thus the principal's action changes the marginal productivity of effort but not the probability of high output without effort. Before accounting for the effort-dependent private benefit under B , effort creates net surplus under A but costs net surplus under B .

Before choosing the action, the principal privately observes $Z \sim F$. Throughout, the support of F is contained in $[0, \infty)$, F is atomless, and $\mathbb{E}Z < \infty$. Results that use derivatives additionally assume that F admits a continuous strictly positive density f on the relevant interior region. If the principal chooses B , she obtains an additional private benefit

$$Z(\kappa + e), \quad \kappa > 0.$$

The term κZ is the component of the private benefit that does not depend on effort. Under B , effort increases the private benefit by Z . Thus κ determines the size of the effort-independent component relative to the effort-dependent component. This captures situations in which employee effort helps management realize the private benefit associated with B . The principal's and agent's payoffs are, respectively,

$$y - w(y) + \mathbf{1}\{a = B\}Z(\kappa + e), \quad w(y) - ce.$$

The same principal chooses the ex ante design and later chooses the action. Thus the private benefit is part of her objective at both stages.

Before observing Z , the principal offers a contract contingent on verifiable output. Since output

⁵Section 7 relaxes the binary structure by considering richer output, richer effort, and more than two managerial actions.

is binary, any such contract is fully described by $w(0)$ and $w(1)$, and can equivalently be written as

$$w(y) = s + by, \quad b \geq 0, \quad (1)$$

where $s = w(0)$ and $b = w(1) - w(0)$ is the output bonus. (1) imposes no restriction on contracts contingent only on output. The restriction $b \geq 0$ is also without loss. A negative bonus never induces effort. If effort is not induced, a constant contract gives the agent the same expected utility and leaves the principal's action choice unchanged. The fixed payment s is unrestricted in the main analysis. The agent's reservation utility is $\bar{u}$. Throughout, we consider the principal's payoff after using the fixed payment to make the participation constraint bind.

The principal also commits to a disclosure policy that generates internal signals about the chosen action. Let Σ be an arbitrary signal space. The policy specifies the conditional distribution

$$\pi(d\sigma | a)$$

of a signal $\sigma \in \Sigma$ after action a is chosen. The policy cannot condition directly on Z . The agent observes σ before choosing effort. The signal is unverifiable and cannot enter the contract in (1). Section 6.2 studies the case in which it can.

The disclosure policy is fixed before Z is observed, while the productive action remains a managerial decision taken after the principal learns Z . The policy can reveal the action but not the privately known benefit that motivated it. The present problem arises because an organization can disclose actions more readily than the motives and opportunities, represented by Z , that shape those actions.

The timing is as follows.

1. The principal chooses the contract and disclosure policy, and the agent decides whether to participate.
2. The principal observes Z and chooses A or B .
3. The internal signal is generated.
4. The agent observes the signal and chooses effort.
5. Output and compensation are realized.

Define

$$g_A \equiv \Delta_A - c > 0, \quad g_B \equiv c - \Delta_B > 0.$$

The term g_A is the net surplus created by effort under A . The term g_B is the net cost of effort under B before accounting for the additional private benefit Z . Effort under B therefore increases total surplus by $Z - g_B$.

An equilibrium consists of a principal action strategy, an agent effort strategy, and beliefs derived from the disclosure policy and the action strategy. At indifference, any mixture between effort and no effort is a best response. When an implementation requires a particular mixture, we take that

mixture as part of the agent's equilibrium strategy. If the principal is indifferent at an action cutoff, either action may be chosen. Since F is atomless, the principal's tie-breaking at the cutoff does not affect equilibrium probabilities or payoffs. The equilibrium probability that the principal chooses A is generated by her action strategy rather than fixed in advance. A given contract and disclosure policy may support multiple equilibria. We use principal-preferred equilibrium selection: the design is evaluated at the equilibrium that gives the principal the highest ex ante payoff. Section 5 explains how partial disclosure can generate such multiplicity.

4 Reducing the disclosure problem

For a signal σ , let

$$\mu(\sigma) = \Pr(A \mid \sigma).$$

The agent's expected increase in the probability of $y = 1$ from the agent's effort is

$$D(\sigma) = \mu(\sigma)\Delta_A + [1 - \mu(\sigma)]\Delta_B.$$

The agent exerts effort if and only if

$$bD(\sigma) \geq c.$$

Since $D(\sigma)$ is the expected effect of the agent's effort on the probability of $y = 1$, $bD(\sigma)$ is the agent's expected increase in bonus payment from effort. The probability of high output without effort, p^0 , drops out. Internal disclosure matters only because the principal's action changes the marginal return to the agent's effort.

4.1 Two-signal reduction

Let us allow the agent to mix after a signal. Let $\rho(\sigma) \in [0, 1]$ be the probability that the agent exerts effort following σ . Define

$$\alpha = \int_{\Sigma} \rho(\sigma)\pi(d\sigma \mid A), \quad \beta = \int_{\Sigma} \rho(\sigma)\pi(d\sigma \mid B). \quad (2)$$

These are the probabilities of the agent exerting effort conditional on the two principal actions.

Lemma 1 (Two internal signals are sufficient). *For every equilibrium under an arbitrary action disclosure policy, there is an outcome-equivalent policy with two internal signals, h and ℓ , after which the agent exerts effort and does not exert effort, respectively. The policy satisfies*

$$\Pr(h \mid A) = \alpha, \quad \Pr(h \mid B) = \beta.$$

This simply groups together signals that induce the same effort choice (Myerson, 1982; Bergemann and Morris, 2016). Crucially, the construction preserves the principal's payoff from each action for

every Z , and hence preserves her action choice. The detailed variation among signals that induce the same effort response is irrelevant.

4.2 One-sided disclosure

Since exerting effort is optimal for the agent after h but not after ℓ , h must make A at least as likely as ℓ does. Thus $\alpha \geq \beta$. Call an equilibrium *nondegenerate* if the agent exerts effort with positive probability and the principal chooses each action with positive probability. The cases in which B is always chosen are considered in Section 5.

A nondegenerate equilibrium must have $b < 1$. If $b \geq 1$, the principal weakly prefers B even at $Z = 0$, so both actions cannot be chosen in such an equilibrium.

Proposition 1 (One-sided disclosure). *Fix a bonus b . Every disclosure policy that is optimal given b and induces a nondegenerate equilibrium has an outcome-equivalent form*

$$\Pr(h \mid A) = 1, \quad \Pr(h \mid B) = q$$

for some $q \in [0, 1]$. The agent exerts effort after h and does not exert effort after ℓ .

We henceforth write $q = \Pr(h \mid B)$. Full disclosure corresponds to $q = 0$: action A generates h and action B generates ℓ . No disclosure corresponds to $q = 1$, in which case ℓ has zero probability. The disclosure problem has therefore been reduced to choosing q .

The proof compares any two-signal policy with one that increases the probability of h after both actions in the same proportion until A always generates h . For a fixed frequency of the two actions, this leaves what h says about the action unchanged while making the agent exert effort more often. At the current cutoff, the principal is indifferent between the two actions. Since B also provides the effort-independent private benefit κZ , the part of the principal's payoff generated through effort must favour A . Making h more frequent magnifies this part of the payoff, so some types near the existing cutoff switch from B to A , and the cutoff rises. The resulting increase in the frequency of A makes h more indicative of A , so exerting effort after h remains optimal for the agent. The proof shows that the principal's expected payoff rises in this comparison. At the endpoint, A always generates h , while ℓ can only follow B , so no effort after ℓ is optimal because $b\Delta_B < c$. The optimal policy therefore always generates h after A and generates it with weakly lower probability after B . Strict partial disclosure corresponds to $q \in (0, 1)$.

5 Joint contract and disclosure design

Proposition 1 reduces disclosure to the probability q that B generates the signal followed by effort. The joint problem then has two layers. For each q , the bonus and the principal's subsequent action cutoff must be consistent with the agent's effort incentive. The principal then chooses q , anticipating this equilibrium response.

5.1 Managerial action and incentives

Under the one-sided policy, action A always generates h , while action B generates h with probability q . For a given bonus b and probability q , the principal obtains

$$(1 - b)\Delta_A$$

from A , excluding terms common across actions. The agent always exerts effort under A , so the principal keeps $1 - b$ of the effect of effort on expected output. On the other hand, from B she obtains

$$\kappa Z + q[(1 - b)\Delta_B + Z].$$

The agent exerts effort under B only when h is generated, which occurs with probability q . When the agent exerts effort, it raises both expected output and the private benefit by Z . She therefore chooses A if and only if

$$Z \leq t(b, q) \equiv \frac{(1 - b)(\Delta_A - q\Delta_B)}{\kappa + q}. \quad (3)$$

Thus $t(b, q)$ is the highest Z for which the principal chooses A . Both a higher q and a higher bonus make B more attractive to the principal:

$$\frac{\partial t}{\partial q} = -\frac{(1 - b)(\Delta_A + \kappa\Delta_B)}{(\kappa + q)^2} < 0, \quad \frac{\partial t}{\partial b} = -\frac{\Delta_A - q\Delta_B}{\kappa + q} < 0. \quad (4)$$

A higher q makes the agent exert effort more often under B , increasing both expected output and the private benefit under that action. A higher bonus also makes B relatively more attractive. Under A , the agent always exerts effort and raises the probability of high output by Δ_A , whereas under B the agent exerts effort only when h is generated, which occurs with probability q , and raises that probability by Δ_B . Increasing the bonus therefore raises the principal's expected bonus payment attributable to effort more under A than under B .

Now fix q and suppose the cutoff is t . Then A is chosen with probability $F(t)$ and always generates h , while B is chosen with probability $1 - F(t)$ and generates h with probability q . Thus A contributes probability mass $F(t)$ to h , while B contributes $q[1 - F(t)]$. Their relative masses determine what the agent learns from h . Conditional on h , the expected effect of effort on the probability of $y = 1$ is

$$D_h(t, q) = \frac{F(t)\Delta_A + q[1 - F(t)]\Delta_B}{F(t) + q[1 - F(t)]}. \quad (5)$$

The agent's expected bonus gain from effort is therefore $bD_h(t, q)$, so he exerts effort after h if and only if

$$bD_h(t, q) \geq c. \quad (6)$$

For $q < 1$, signal ℓ is generated only after B . Since $b < 1$ in every nondegenerate partial-disclosure solution and $\Delta_B < c$, no effort after ℓ is optimal. When $q = 1$, signal ℓ has zero probability.

For a given disclosure policy, the bonus required to induce effort depends on how likely the

principal is to choose A , because that probability determines what the effort-inducing signal conveys about the action. But the probability that the principal chooses A itself depends on the bonus, since the bonus changes her relative payoff from A and B . The optimal contract must therefore satisfy a fixed-point condition: its bonus must be exactly the bonus required given the probability that the principal chooses A under that contract.

The degree of disclosure determines the strength of this feedback. Under full disclosure, the effort-inducing signal identifies A , so the required bonus does not depend on how often A is chosen. As more B choices are pooled into that signal, what the signal conveys increasingly depends on the probability that the principal chooses A . More pooling makes the signal less indicative of A and raises the required bonus. The higher bonus in turn makes A less attractive to the principal, reducing the probability that she chooses A and further weakening what the signal conveys about the return to effort.

At a nondegenerate optimum, the incentive condition binds, as shown in Section 5.2. Rearranging the principal's cutoff condition and imposing the binding incentive condition give

$$b = 1 - \frac{(\kappa + q)t}{\Delta_A - q\Delta_B}, \quad (7)$$

$$b = \frac{c}{D_h(t, q)}. \quad (8)$$

The first expression gives the bonus that generates cutoff t . The second gives the smallest bonus required when the cutoff is t . The fixed point requires these two calculations to give the same bonus.

There are two distinct sources of multiplicity. First, for a given q , the two conditions can have more than one intersection because the required-bonus relation depends on $F(t)$, which determines how a change in the cutoff changes what the agent learns from h . These intersections correspond to different bonuses, not to different positive-effort cutoffs under the same bonus. For any given (b, q) , conditional on the agent exerting effort after h , the principal's payoff difference between A and B is strictly decreasing in Z , so her cutoff is unique. As shown below, for each q the principal selects the feasible intersection with the largest cutoff.

Second, partial disclosure creates a separate source of equilibrium multiplicity for a fixed contract. When $q > 0$ and $b\Delta_B < c$, there is an all- B , no-effort equilibrium. If the agent expects not to exert effort after either signal, B gives every principal type with $Z > 0$ the additional private benefit κZ , so all such types choose B . Every signal observed on path then indicates B , which validates the agent's decision not to exert effort. In the relevant $b < 1$ region, if even the candidate outcome in which the agent always exerts effort after h satisfies

$$bD_h(t(b, q), q) < c,$$

then the all- B , no-effort outcome is the unique equilibrium.⁶ Economically, the no-effort outcome

⁶If effort follows h with probability x , both the induced cutoff and the agent's return to effort after h increase with x . Failure of the incentive condition at $x = 1$ therefore rules out every $x > 0$.

is harder to escape when h is generated more often by B or when B 's private benefit is stronger, since either force makes h less favourable evidence about A . In the nondegenerate partial-disclosure region, $b < 1$ and $\Delta_B < c$, so $b\Delta_B < c$ automatically. Thus whenever a positive-effort equilibrium exists under partial disclosure, the all- B , no-effort equilibrium also exists. As specified in Section 3, we evaluate a design at the equilibrium preferred by the principal. At the full-disclosure contract considered below, by contrast, h can only follow A , so the agent's response to h does not depend on how frequently A is chosen.

Lemma 2 (Bonus response). *Suppose that along a nondegenerate local solution the smallest bonus satisfying*

$$bD_h(t(b, q), q) = c$$

can be written as a differentiable function $b(q)$, with $t(q) \equiv t(b(q), q)$, and that

$$\frac{\partial}{\partial b} [bD_h(t(b, q), q)] > 0.$$

Then the bonus required to induce effort after h rises with the probability that B generates h :

$$b'(q) > 0.$$

At $q = 0$, $b'(q)$ is understood as a right derivative.

Under the same conditions, the induced cutoff also falls. Since both partial derivatives in (4) are negative and $b'(q) > 0$,

$$t'(q) = \frac{\partial t}{\partial q} + \frac{\partial t}{\partial b} b'(q) < 0. \quad (9)$$

At $q = 0$, the derivative is understood as a right derivative. The first term in (9) is the direct effect of greater pooling at a fixed bonus. The second is the additional effect through the higher bonus. Thus joint contract design amplifies the fall in the cutoff caused by greater pooling.

For the analysis below, it is convenient to eliminate the bonus from the two fixed-point conditions. For $q > 0$, combining (7) and (8) gives

$$F(t) = \frac{q [g_B(\Delta_A - q\Delta_B) + \Delta_B t(\kappa + q)]}{(\Delta_A - q\Delta_B) [g_A + qg_B - t(\kappa + q)]}. \quad (10)$$

This is a single equation for a cutoff satisfying both conditions. It is used only for $q > 0$.

5.2 Optimal disclosure

We now choose q , taking into account the cutoff and bonus generated by the equilibrium described above. To express the principal's expected payoff, define

$$M(t) = \int_t^\infty z \, dF(z), \quad H(t) = \int_t^\infty (z - g_B) \, dF(z).$$

Types above t choose B . When B is chosen with positive probability,

$$M(t) = [1 - F(t)]\mathbb{E}[Z \mid B], \quad H(t) = [1 - F(t)]\mathbb{E}[Z - g_B \mid B].$$

Thus $\kappa M(t)$ is the principal's expected effort-independent private benefit from choosing B , while $qH(t)$ is the expected change in total surplus from the agent exerting effort when B is chosen and generates h .

With an unrestricted fixed payment, the participation constraint binds. Holding the induced effort and action choices fixed, a higher expected bonus payment can be offset by a lower fixed payment while leaving the agent at his reservation utility. The bonus is therefore not a direct resource cost. It matters for the principal's expected payoff because it changes the agent's effort and the principal's managerial action choice. After using the fixed payment to satisfy the participation constraint, the principal's expected payoff is

$$\mathcal{W}(t, q) = p^0 - \bar{u} + g_A F(t) + \kappa M(t) + qH(t).$$

The term $g_A F(t)$ is the surplus generated by effort under A . The other two terms have the interpretations given above.

To understand the cost of changing the cutoff, first suppose that for a fixed q the cutoff could be chosen directly to maximize total surplus. For a marginal type $Z = t$, choosing A generates net surplus g_A from effort, excluding the common term p^0 . Choosing B generates the effort-independent private benefit κt and, with probability q , the additional net surplus $t - g_B$ from effort. If A generated more surplus than B at the marginal type, raising the cutoff slightly would move nearby types from B to A and increase total surplus. If B generated more, lowering the cutoff would increase total surplus. The surplus-maximizing cutoff therefore satisfies the explicit indifference condition

$$g_A = \kappa t^S(q) + q[t^S(q) - g_B], \quad (11)$$

so

$$t^S(q) = \frac{g_A + qg_B}{\kappa + q}.$$

At $t^S(q)$, the two actions generate the same total surplus for the marginal type. Equivalently,

$$\frac{\partial \mathcal{W}(t, q)}{\partial t} = (\kappa + q)f(t)[t^S(q) - t]. \quad (12)$$

The incentive condition limits how high the cutoff can be.

Lemma 3 (Feasible cutoffs). *For every $q > 0$, any feasible (b, q, t) that induces effort after h and uses both actions has*

$$t < t^S(q).$$

For $q = 0$, the largest feasible cutoff is $t^S(0) = g_A/\kappa$. Consequently, $\mathcal{W}(t, q)$ is weakly increasing in t throughout the feasible set for every q , and strictly increasing while both actions are chosen.

For $q > 0$, the lemma implies that every feasible cutoff lies below the surplus-maximizing cutoff. Hence some types choose B even though choosing A for them would generate higher total surplus. This is why a disclosure-induced fall in the cutoff reduces total surplus.

For any q that supports a feasible cutoff, let $\tau(q)$ denote the largest feasible cutoff, meaning the largest $t \geq 0$ for which the bonus in (7) lies in $[0, 1]$ and the agent's incentive condition holds. Let

$$\mathcal{Q} = \{q \in [0, 1] : \text{a feasible cutoff exists}\}.$$

Thus $\tau(q)$ is the highest managerial cutoff attainable at pooling probability q , and $\mathcal{Q}$ is the set of pooling probabilities that support such an outcome. If F has bounded support, an all- A outcome with effort is already included whenever $F(t) = 1$. Lemma 3 implies that the principal selects $\tau(q)$. It also implies that the agent's incentive condition binds at any nondegenerate optimum. Otherwise the bonus could be reduced, raising the cutoff and the principal's payoff while preserving effort after h .

The remaining candidate outcomes are simple. Under full disclosure, $q = 0$, the smallest bonus inducing effort under A is $b = c/\Delta_A$. The associated cutoff and payoff are

$$t_{\text{FD}} = \frac{g_A}{\kappa}, \quad V_{\text{FD}} = p^0 - \bar{u} + g_A F(t_{\text{FD}}) + \kappa M(t_{\text{FD}}). \quad (13)$$

Full disclosure can instead induce effort under both actions. In the present parameter region this requires $b \geq c/\Delta_B > 1$ and makes the principal choose B for every positive Z . Its payoff is

$$V_{B,1} = p^0 - \bar{u} + (1 + \kappa)\mathbb{E}Z - g_B. \quad (14)$$

If the agent never exerts effort, the principal also chooses B for every positive Z , giving

$$V_{B,0} = p^0 - \bar{u} + \kappa\mathbb{E}Z. \quad (15)$$

No disclosure, $q = 1$, is included in the reduced problem whenever the agent is willing to exert effort after the uninformative signal. Otherwise the best all- B outcome with effort is (14).

Proposition 2 (Joint contract and disclosure design). *The joint contract and disclosure problem has a solution. Every nondegenerate solution with $q > 0$ is represented by a pair (t, q) satisfying (10), with the bonus given by (8). For each $q \in \mathcal{Q}$, the principal selects the largest feasible cutoff $\tau(q)$. Hence the nondegenerate problem with effort after h reduces to*

$$\max_{q \in \mathcal{Q}} \mathcal{W}(\tau(q), q). \quad (16)$$

At $q = 0$, the payoff is V_{FD} in (13). The global optimum is the maximum of the reduced payoff in (16) and the candidate payoffs $V_{B,1}$ and $V_{B,0}$ in (14) and (15).

Strict partial disclosure trades off the additional surplus from inducing effort after some B choices against the surplus loss caused by the lower cutoff. As (9) shows, more pooling lowers the cutoff directly and also raises the required bonus, which lowers the cutoff further. It can make productive effort possible under B , but it also expands the range of Z for which the principal chooses B even though A would generate greater total surplus.

5.3 A partial disclosure example

Consider

$$\Delta_A = 0.9, \quad \Delta_B = 0.2, \quad c = 0.5, \quad \kappa = 0.4, \quad (17)$$

with $p^0 = \bar{u} = 0$ and Z exponentially distributed with mean 0.5.

Under full disclosure, the smallest bonus is $b = c/\Delta_A = 5/9$. The agent exerts effort after A and not after B . The principal chooses A when

$$Z \leq t_{\text{FD}} = \frac{\Delta_A - c}{\kappa} = 1.$$

Thus A is chosen with probability 0.8647 and B with probability 0.1353. Full disclosure does not remove B from the principal's choice set. Its payoff is

$$V_{\text{FD}} = 0.4271.$$

The optimal partial disclosure policy is approximately

$$q^* = 0.2946, \quad b^* = 0.6446, \quad t^* = 0.4304.$$

The principal chooses A with probability 0.5772. Following signal h , the posterior probability of A is 0.8225, and the expected effect of effort on the probability of $y = 1$ is 0.7757. Hence $b^* D_h(t^*, q^*) = 0.5 = c$. The resulting payoff is

$$V^* = 0.4668.$$

The best no-disclosure outcome induces effort for every realization of Z and gives payoff $V_{B,1} = 0.4$.

Thus

$$V^* > V_{\text{FD}} > V_{B,1}.$$

Table 1 collects the three disclosure policies.

Table 1: The three disclosure policies in the numerical example

	q	b	$\Pr(A)$	Payoff
Full disclosure	0	0.5556	0.8647	0.4271
Optimal partial disclosure	0.2946	0.6446	0.5772	0.4668
Best no disclosure	1	2.5	0	0.4

Notes. The parameters are given in (17). Under no disclosure, the best outcome induces effort for every realization of Z . The bonus shown is the smallest bonus, c/Δ_B , that supports it.

Figure 1 plots the principal's payoff under partial disclosure and the full-disclosure and no-disclosure payoffs against q . The partial-disclosure curve ends at approximately $q = 0.3619$. Beyond that point no cutoff satisfies the principal's cutoff condition and the agent's incentive condition simultaneously, so no nondegenerate solution with effort after h remains feasible. Economically, as more B is pooled into h , the signal becomes less favourable news about the return to effort and the required bonus rises. Greater pooling lowers the cutoff directly, and the higher bonus lowers it further. Eventually no bonus can both induce the agent to exert effort after h and sustain an equilibrium in which both actions are chosen.

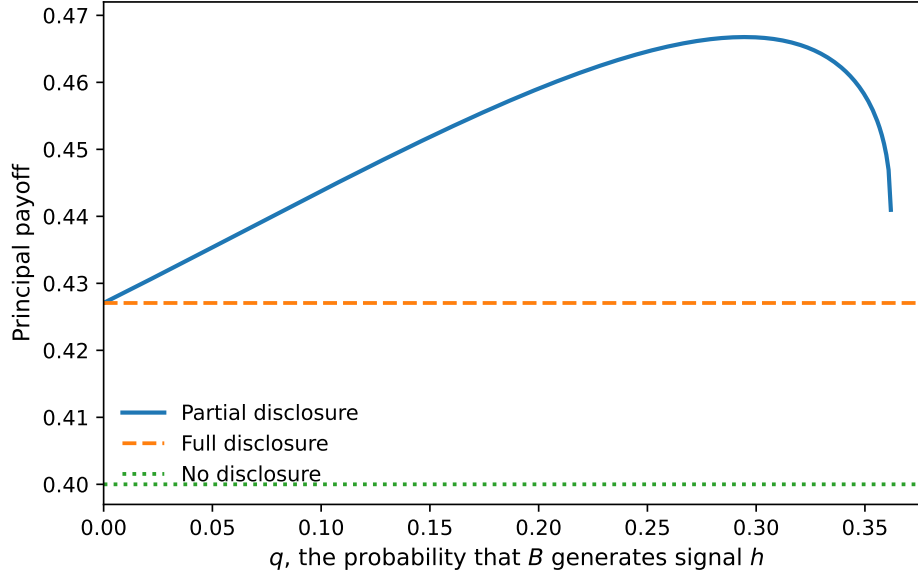


Figure 1: Partial disclosure, full disclosure, and no disclosure

Notes. The parameters are given in (17). For each q , the solid line gives the principal's payoff at the largest feasible cutoff under partial disclosure. It ends at approximately $q = 0.3619$, where the two feasible fixed-point intersections merge and no nondegenerate feasible solution remains. The sharp endpoint is therefore not a discontinuity. The dashed line is the best full-disclosure payoff. The dotted line is the best no-disclosure payoff.

The example shows that the result does not rely on making B unavailable or unproductive under full disclosure. Under full disclosure, the principal chooses B for large realizations of Z even though the agent does not exert effort. Partial disclosure allows the agent to exert effort after some B choices, creating additional surplus when Z is sufficiently large. Consistent with (9), pooling lowers the cutoff directly, and the increase in the required bonus from 0.5556 under full disclosure to 0.6446 at the partial-disclosure optimum lowers it further. The cutoff therefore falls from 1 to 0.4304. Thus B is chosen whenever $Z > 0.4304$ rather than only when $Z > 1$. At q^* , the surplus-maximizing cutoff is $t^S(q^*) = 0.7031$, so the actual cutoff 0.4304 means that the principal chooses B for some realizations of Z even though A would generate greater total surplus. The optimum balances this loss against the additional surplus generated by effort under B . This tradeoff makes strict partial disclosure optimal.

6 Sources of partial disclosure

The partial-disclosure result in the main model relies on two features. The principal cannot commit to how her action depends on Z , and the internal signal cannot enter the contract. This section relaxes each restriction in turn.

6.1 Commitment to the managerial action

Suppose that before Z is observed the principal can commit to a rule specifying the action for each Z . For any such rule, define

$$\lambda = \Pr(A), \quad m_B = \mathbb{E}[Z \mathbf{1}\{B\}].$$

Here λ is the probability that the committed rule selects A , while m_B is the probability-weighted mean of Z over realizations assigned to B . For a cutoff rule that assigns B when $Z > t$, $m_B = M(t)$. This quantity remains well defined when one action is chosen with probability zero. Since disclosure can condition only on the principal's action, the probability that the agent exerts effort following B cannot vary with Z . Let $x_A, x_B \in [0, 1]$ denote the probabilities that the agent exerts effort following A and B , respectively, including after an off-path action. The principal's ex ante payoff conditional on the committed rule is

$$\begin{aligned} V^C(x_A, x_B) = & p^0 - \bar{u} + \kappa m_B + \lambda x_A g_A \\ & + x_B [m_B - (1 - \lambda) g_B]. \end{aligned} \tag{18}$$

Proposition 3 (Commitment to the managerial action). *Fix any such committed rule. If the principal jointly optimizes the common output-contingent contract and disclosure policy, full action disclosure is weakly optimal conditional on that rule. The payoff-maximizing values of x_A and x_B can be chosen as*

$$x_A^* = 1, \quad x_B^* = \mathbf{1}\{m_B - (1 - \lambda)g_B \geq 0\}. \quad (19)$$

When B is chosen with positive probability, the condition for $x_B^ = 1$ is equivalent to $\mathbb{E}[Z \mid B] + \Delta_B \geq c$. Consequently, full action disclosure is weakly optimal in the joint problem over how the action depends on Z , the contract, and the disclosure policy.*

Full disclosure need not be uniquely optimal. If the agent should exert effort after both actions, no disclosure can implement the same effort choices. Once the action rule is fixed, however, disclosure no longer changes which action the principal chooses. Full disclosure can implement the payoff-maximizing effort choice following each action, so pooling cannot strictly improve the principal's payoff by changing managerial action choice. Implementing effort after B may require $b \geq c/\Delta_B > 1$. This is feasible here because the principal has committed to her action rule, so the bonus no longer affects which managerial action she chooses.

6.2 Contractible signals

Now return to discretionary action choice but allow compensation to depend on the internal signal. Consider an arbitrary disclosure policy with contractible signals. For each action a , let $x_a \in [0, 1]$ denote the probability that the agent exerts effort following that action and let T_a denote the corresponding expected transfer, including for an off-path action. These two objects summarize the agent's effort choice and expected compensation following each action. For type Z , the payoffs are

$$p^0 + x_A \Delta_A - T_A \quad (20)$$

from A , and

$$p^0 + x_B \Delta_B - T_B + \kappa Z + x_B Z \quad (21)$$

from B .

Proposition 4 (Contractible signals). *Every outcome attainable when the internal signal is contractible has an outcome-equivalent implementation under full action disclosure and contracts contingent on both the principal's action and output. Hence full action disclosure is without loss when the internal signal can enter the contract.*

For each action a , the proof constructs the output-contingent contract

$$w_a(y) = s_a + b_a y, \quad b_a = \frac{c}{\Delta_a}, \quad s_a = T_a - \frac{c}{\Delta_a}(p^0 + x_a \Delta_a). \quad (22)$$

Since output is binary, this describes an unrestricted contract conditional on the realized action and output. The agent is indifferent, so exerting effort with probability x_a is a best response under

the convention stated in Section 3. Equivalently, the full-disclosure implementation can add an independent public lottery while continuing to reveal the action. The expected transfer and the probability that the agent exerts effort conditional on each action are preserved, so the principal's payoff from each action is preserved for every Z .

The proposition is specific to the present binary-effort model. When the signal is contractible, full disclosure makes the realized action contractible as well. The principal can then use action-specific contracts to reproduce the probability that the agent exerts effort and the expected transfer associated with each action. Pooling is no longer needed to create different incentives through a common bonus, so full disclosure is without loss.

7 Extensions and robustness

This section relaxes binary output, binary effort, the two-action assumption, and the unrestricted fixed payment. The first two extensions distinguish the two-signal and one-sided results from the feedback between disclosure, effort, and managerial action choice.

7.1 Richer output

Let output Y be integrable and let its distribution depend on the managerial action and effort. Assume that the no-effort distribution is common across actions. Write an arbitrary output-contingent contract as $w(Y) = s + v(Y)$, where v is an integrable function of output. Fix v and let s adjust so that the participation constraint binds. For each action $a \in \{A, B\}$, define

$$\begin{aligned}\gamma_a &= \mathbb{E}[v(Y) \mid a, e = 1] - \mathbb{E}[v(Y) \mid a, e = 0], \\ \chi_a &= \mathbb{E}[Y - v(Y) \mid a, e = 1] - \mathbb{E}[Y - v(Y) \mid a, e = 0].\end{aligned}$$

The term γ_a is the increase in the agent's expected compensation from effort under action a , while χ_a is the increase in expected output net of $v(Y)$ that effort creates for the principal. After a signal that gives posterior probability μ to action A , the agent exerts effort if and only if

$$\mu\gamma_A + (1 - \mu)\gamma_B \geq c.$$

Proposition 5 (Richer output). *For every equilibrium under an arbitrary disclosure policy, there is an outcome-equivalent policy with two signals, one followed by effort and the other by no effort. Suppose in addition that*

$$\gamma_A \geq c > \gamma_B, \quad \chi_A > \chi_B \geq 0.$$

Then every optimal disclosure policy for that v that induces a nondegenerate equilibrium has an outcome-equivalent form

$$\Pr(h \mid A) = 1, \quad \Pr(h \mid B) = q$$

for some $q \in [0, 1]$, with effort after h and no effort after ℓ .

The two-signal result uses binary effort, not binary output. The first ordering, $\gamma_A \geq c > \gamma_B$, says that the agent has a stronger reason to exert effort under A than under B . The second, $\chi_A > \chi_B \geq 0$, says that the principal also gains more from effort under A . These are exactly the two comparisons used by the one-sided argument. Starting from a two-signal policy, increase the probability of h after both actions in the same proportion. At the marginal realization of Z , the effort-dependent part of the principal's payoff must favour A because B also gives her the effort-independent private benefit κZ . Making h more frequent therefore makes A relatively more attractive and raises the cutoff. As A is chosen more often, h becomes better news about the agent's return to effort, so the agent continues to exert effort after h . In the binary-output model, $\gamma_a = b\Delta_a$ and $\chi_a = (1 - b)\Delta_a$, so these conditions hold in every nondegenerate equilibrium. Binary output is therefore not essential to the two-signal and one-sided results.

7.2 Richer effort

Return to binary output and let the agent choose effort from a set $\mathcal{E} \subseteq [0, \bar{e}]$, with cost $C(e)$ and $C(0) = 0$. Assume $p^0 + \bar{e}\Delta_A \leq 1$. After a signal with posterior probability μ of action A , the part of the agent's payoff that depends on effort is

$$be[\mu\Delta_A + (1 - \mu)\Delta_B] - C(e).$$

Lemma 4 (Finite effort levels). *If $\mathcal{E} = \{e_1, \dots, e_K\}$, every equilibrium under an arbitrary disclosure policy has an outcome-equivalent form with at most K signals, one for each effort level used with positive probability. This form preserves the conditional distribution of effort given each managerial action and preserves the principal's action choice for every Z .*

For binary effort this is the two-signal reduction in Lemma 1. With several effort levels, different signals can induce different levels of effort, so the exact two-signal and one-sided results need not survive.

Now let $\mathcal{E} = [0, \bar{e}]$ and suppose C is increasing, continuous, and strictly convex. The agent's

effort response is unique and can be written

$$e(\mu; b) = \arg \max_{e \in [0, \bar{e}]} \{be[\mu\Delta_A + (1 - \mu)\Delta_B] - C(e)\}.$$

Fix a bonus b and disclosure policy π . Suppose the principal chooses A for $Z \leq t$. Bayes' rule then gives a posterior $\mu_t(\sigma)$ after each signal. Define expected effort conditional on each action by

$$\bar{e}_a(t) = \int e(\mu_t(\sigma); b)\pi(d\sigma | a), \quad a \in \{A, B\},$$

where the dependence on b and π is suppressed.

Proposition 6 (Richer effort). *For a fixed bonus and disclosure policy, an interior cutoff t is consistent with the principal's action choice if and only if*

$$t = \frac{(1 - b)[\Delta_A \bar{e}_A(t) - \Delta_B \bar{e}_B(t)]}{\kappa + \bar{e}_B(t)}. \quad (23)$$

In (23), the numerator is the difference between the principal's expected output net of bonus payments attributable to effort under A and under B . The denominator is the private-benefit advantage of B per unit of Z , including the part generated by expected effort under B . The same fixed-point logic therefore survives. The cutoff determines the frequencies of the two actions. Those frequencies determine what the agent infers from pooled signals, which determines the agent's effort. Expected effort under each action then changes the principal's relative payoff from A and B , and hence the cutoff.

Under full disclosure, the agent knows the action, so what he learns does not depend on the cutoff. Whenever a signal pools the two actions, the probability he assigns to A generally depends on $F(t)$. For example, under the one-sided two-signal policy,

$$\mu_h(t, q) = \frac{F(t)}{F(t) + q[1 - F(t)]}.$$

Writing $e_h = e(\mu_h(t, q); b)$ and $e_\ell = e(0; b)$ gives

$$\bar{e}_A(t) = e_h, \quad \bar{e}_B(t) = qe_h + (1 - q)e_\ell.$$

With binary effort, $e_h = 1$ and $e_\ell = 0$, and (23) reduces to (3). Richer effort can require more signals. What survives is the same feedback: pooling makes what the agent learns depend on the cutoff, what he learns changes his effort, and expected effort changes the cutoff.

7.3 More managerial actions

Return to binary effort and output, and let the principal choose $i \in \{1, \dots, n\}$. Under action i , $y = 1$ occurs with probability $p^0 + e\Delta_i$, so Δ_i is the increase in the probability of $y = 1$ when the

agent exerts effort under that action. The principal's nonoutput payoff from action i is

$$\phi_i(Z) + e\psi_i(Z).$$

Here $\phi_i(Z)$ is the component independent of effort, while $\psi_i(Z)$ is the additional nonoutput payoff generated by effort. Fix the common bonus b and define

$$G_i = b\Delta_i - c, \quad R_i(Z) = (1 - b)\Delta_i + \psi_i(Z).$$

Here G_i is the agent's net gain from effort and $R_i(Z)$ is the principal's gain from effort. Their sum is the total surplus created by effort under action i .

The argument behind Lemma 4 does not depend on the number of managerial actions. With binary effort, two signals therefore suffice. Let $q_i = \Pr(h \mid i)$, where the agent exerts effort after h , and write $\mathbf{q} = (q_1, \dots, q_n)$. Thus q_i is the probability that action i generates the signal followed by effort. The principal chooses the action that maximizes $\phi_i(Z) + q_i R_i(Z)$. Let $Z_i(\mathbf{q})$ be the set of realizations of Z for which action i is optimal given $\mathbf{q}$. After using the fixed payment to satisfy the participation constraint, the principal's expected payoff is

$$\mathcal{V}(\mathbf{q}) = p^0 - \bar{u} + \sum_{i=1}^n \int_{Z_i(\mathbf{q})} [\phi_i(z) + q_i(R_i(z) + G_i)] dF(z).$$

The agent's probability-weighted net gain from effort after h is

$$I(\mathbf{q}) = \sum_{i=1}^n \int_{Z_i(\mathbf{q})} q_i G_i dF(z). \quad (24)$$

Whenever h occurs with positive probability, $I(\mathbf{q})$ equals $\Pr(h)$ times the agent's net gain from effort conditional on h . Thus effort after h is optimal if and only if $I(\mathbf{q}) \geq 0$.

Proposition 7 (More managerial actions). *Suppose the principal's best action is unique for almost every Z . At any point where $\mathcal{V}$ and I are differentiable,*

$$\frac{\partial \mathcal{V}}{\partial q_k} = \int_{Z_k(\mathbf{q})} R_k(z) dF(z) + \frac{\partial I}{\partial q_k}. \quad (25)$$

The first term is the principal's direct gain from inducing effort more often under action k . The second is the total change in the agent's probability-weighted net gain from effort after h . It includes both the direct effect of changing q_k and the change produced when the principal switches actions for some realizations of Z . Thus increasing the probability that action k generates h has a direct payoff effect, but it can also tighten the agent's incentive condition because the induced changes in managerial action choice alter what h conveys. The identity does not require the actions to be ordered by a single cutoff. With two actions, these changes are summarized by a single cutoff. With many actions, changing disclosure can reallocate different realizations of Z among several managerial actions. The same feedback remains: managerial choice changes the composition of the

signal followed by effort, and that composition determines whether the agent is willing to exert effort.

7.4 Limited liability

The main model permits the fixed payment to be negative. Since $b \geq 0$, limited liability is equivalent to $s \geq 0$. At a nondegenerate optimum the agent's effort incentive condition binds, so his expected utility is $s + bp^0$. In the unrestricted problem, the binding participation constraint therefore gives

$$s = \bar{u} - bp^0.$$

If $p^0 = 0$, this payment already satisfies limited liability whenever $\bar{u} \geq 0$. If $p^0 > 0$ and $\bar{u} = 0$, limited liability sets $s = 0$ and leaves the agent rent bp^0 .

Proposition 8 (Limited liability). *If $p^0 = 0$ and $\bar{u} \geq 0$, every partial-disclosure optimum of the unrestricted model with a binding effort incentive condition satisfies limited liability, with $s = \bar{u}$. If $\bar{u} = 0$, limited liability binds exactly.*

Suppose further that at $p^0 = 0$ and $\bar{u} = 0$ the unrestricted problem has a strict global optimum (b^, q^*, t^*) with $0 < q^* < 1$ and both actions chosen. Then for all sufficiently small $p^0 > 0$, the limited-liability problem has an interior partial-disclosure optimum. Limited liability binds, $s = 0$, and the agent obtains positive rent bp^0 .*

Limited liability therefore does not eliminate strict partial disclosure. Along a differentiable interior segment of the selected branch, write $b(q)$ for the bonus associated with $\tau(q)$. When limited liability binds, the principal's expected payoff is $\mathcal{W}(\tau(q), q) - p^0 b(q)$. Lemma 2 implies

$$\frac{d}{dq} [\mathcal{W}(\tau(q), q) - p^0 b(q)] = \frac{d}{dq} \mathcal{W}(\tau(q), q) - p^0 b'(q) < \frac{d}{dq} \mathcal{W}(\tau(q), q).$$

The rent arises because high output occurs with probability $p^0 > 0$ even when the agent does not exert effort. The bonus is therefore sometimes paid without effort. With unrestricted transfers, the fixed payment falls to offset this expected payment, but limited liability prevents that adjustment once $s = 0$. A larger bonus therefore creates additional rent bp^0 . Since more pooling requires a larger bonus along the interior solution, limited liability adds a force against pooling. Proposition 8 also shows that partial disclosure survives for small positive p^0 .

8 Conclusion

This paper studies disclosure of the principal's own productive managerial action. The principal commits to the contract and disclosure policy before learning her private benefit, but chooses the action afterwards. The cutoff governing that choice determines the relative frequencies of A and B behind a pooled signal, and therefore what the agent learns from it. What the signal means

determines the bonus required to induce effort. The bonus in turn changes expected bonus payments across the two actions and hence the cutoff. This fixed point arises because the object being disclosed is generated by a managerial choice that responds to the contract.

In the binary model, every disclosure policy can be reduced to two internal signals, and every optimal policy inducing a nondegenerate equilibrium has a one-sided form. The action for which the agent's effort has the larger effect on expected output always generates the signal followed by effort, while the other generates the same signal with weakly lower probability. Partial disclosure can be uniquely optimal. Pooling B with A can induce the agent to exert effort after some B choices, but it also lowers the cutoff directly and can require a larger bonus, which lowers the cutoff further. As a result, B is chosen for a wider range of Z . Partial disclosure balances the additional surplus generated by effort under B against the surplus loss from the lower cutoff. If the principal can commit to the managerial action and jointly optimize the contract, full disclosure is weakly optimal. Full disclosure is also without loss when the internal signal can enter the contract.

The sharp two-signal and one-sided results use the binary structure, but the feedback between disclosure, the agent's effort, and managerial action choice is more general. Richer output preserves the sharp disclosure results under natural conditions, while richer effort may require more signals without removing the feedback. With many managerial actions, the same tension remains between inducing effort and changing the principal's action choice. Limited liability adds a cost to the larger bonuses associated with pooling but does not eliminate strict partial disclosure. Corporate sales, outsourcing, relocation, restructuring, and technological transitions often leave management with discretion while employee cooperation remains valuable. In such settings, internal communication can change the action management prefers, not only employee reactions to a decision already made. The results explain why disclosure of managerial actions may be partial rather than uniformly full or absent.

A Proofs

A.1 Proof of Lemma 1

Fix an equilibrium under an arbitrary disclosure policy and let P be the unconditional distribution of signals. Let

$$\mu(\sigma) = \Pr(A \mid \sigma), \quad D(\sigma) = \mu(\sigma)\Delta_A + [1 - \mu(\sigma)]\Delta_B.$$

The gain from effort following σ is

$$\Gamma(\sigma) = bD(\sigma) - c.$$

Optimality of the agent's effort choice implies

$$\rho(\sigma) > 0 \implies \Gamma(\sigma) \geq 0, \quad \rho(\sigma) < 1 \implies \Gamma(\sigma) \leq 0.$$

Construct the two-signal policy using (2). Let $\tilde{D}_h$ and $\tilde{D}_\ell$ denote the expected effect of the agent's effort on the probability of $y = 1$ after the two signals. Then

$$\Pr(h)\tilde{D}_h = \int_{\Sigma} \rho(\sigma)D(\sigma) dP(\sigma).$$

Therefore

$$\Pr(h)[b\tilde{D}_h - c] = \int_{\Sigma} \rho(\sigma)\Gamma(\sigma) dP(\sigma) \geq 0.$$

Effort after h is optimal whenever h has positive probability. Similarly,

$$\Pr(\ell)[b\tilde{D}_\ell - c] = \int_{\Sigma} [1 - \rho(\sigma)]\Gamma(\sigma) dP(\sigma) \leq 0,$$

so no effort after ℓ is optimal.

The probability that the agent exerts effort conditional on each action is unchanged. Conditional on A , the principal's payoff relevant for action choice is $\alpha(1 - b)\Delta_A$. Conditional on B , it is

$$\kappa Z + \beta[(1 - b)\Delta_B + Z].$$

These are the original conditional action payoffs. Hence the principal's action choice is preserved for every Z . Expected output, compensation, effort, agent utility, and principal payoff are also preserved.

A.2 Proof of Proposition 1

Consider a two-signal policy from Lemma 1 in which the agent exerts effort only after h , with $\alpha = \Pr(h | A) > 0$ and $\beta = \Pr(h | B)$. Optimality of effort after h requires $\alpha \geq \beta$. Write

$$q = \frac{\beta}{\alpha} \in [0, 1],$$

fix q , and set $\beta = \alpha q$. For $b < 1$, the principal chooses A if and only if

$$Z \leq t(\alpha, q, b) = \frac{\alpha(1 - b)(\Delta_A - q\Delta_B)}{\kappa + \alpha q}.$$

The cutoff rises with α :

$$\frac{\partial t}{\partial \alpha} = \frac{(1 - b)(\Delta_A - q\Delta_B)\kappa}{(\kappa + \alpha q)^2} > 0.$$

Increasing α with q fixed makes h more likely after both actions in the same proportion. It raises the cutoff and therefore raises the posterior probability of A after h . Effort after h remains optimal at every larger α .

Using $M(t)$ and $H(t)$ defined in Section 5.2, the principal's payoff when effort follows h and not ℓ is

$$V(\alpha, q, b) = p^0 - \bar{u} + \kappa M(t) + \alpha[g_A F(t) + qH(t)], \quad (\text{A.1})$$

where $t = t(\alpha, q, b)$. The intermediate values of α are used only to compare the original policy with the endpoint $\alpha = 1$. They need not make no effort optimal after ℓ and need not themselves constitute equilibria. They are used only to establish that the principal's payoff rises along the comparison, while feasibility is required at the endpoint $\alpha = 1$. Differentiate (A.1):

$$V_\alpha = g_A F(t) + qH(t) + f(t)t_\alpha \alpha [b\Delta_A - c + q(c - b\Delta_B)].$$

The second line is nonnegative because $t_\alpha > 0$, $b\Delta_A \geq c$, and $b\Delta_B < c$.

Effort after h is optimal if and only if

$$F(t)(b\Delta_A - c) + q[1 - F(t)](b\Delta_B - c) \geq 0.$$

Adding

$$(1 - b)[F(t)\Delta_A + q[1 - F(t)]\Delta_B] > 0$$

gives

$$F(t)g_A + q[1 - F(t)](\Delta_B - c) > 0.$$

Since

$$qH(t) = q[1 - F(t)](\Delta_B - c) + q \int_t^\infty z dF(z),$$

we obtain $g_A F(t) + qH(t) > 0$. Thus $V_\alpha > 0$. The constructed positive-effort equilibrium at the endpoint $\alpha = 1$ gives strictly higher payoff unless the original policy already has $\alpha = 1$. Under principal-preferred selection, the payoff from the endpoint policy is at least this value. At that endpoint, ℓ is generated only by B , so no effort after ℓ is optimal because $b\Delta_B < c$.

A.3 Proof of Lemma 2

For this proof, let

$$\delta \equiv \Delta_A - \Delta_B, \quad r(t, q) \equiv F(t) + q[1 - F(t)].$$

Direct differentiation of (5) gives

$$\frac{\partial D_h}{\partial q} = -\frac{\delta F(t)[1 - F(t)]}{r(t, q)^2}, \quad \frac{\partial D_h}{\partial t} = \frac{\delta q f(t)}{r(t, q)^2}.$$

Using (4), the change in $D_h(t(b, q), q)$ with respect to q at a fixed bonus is

$$\frac{\partial D_h}{\partial q} + \frac{\partial D_h}{\partial t} \frac{\partial t}{\partial q} = -\frac{\delta F(t)[1 - F(t)]}{r(t, q)^2} - \frac{\delta q(1 - b)(\Delta_A + \kappa \Delta_B)f(t)}{(\kappa + q)^2 r(t, q)^2} < 0.$$

At $q = 0$, this is the right derivative. Differentiating the binding condition $bD_h(t(b, q), q) = c$ along the local solution gives

$$b'(q) = -\frac{b \left[\frac{\partial D_h}{\partial q} + \frac{\partial D_h}{\partial t} \frac{\partial t}{\partial q} \right]}{\frac{\partial}{\partial b} [bD_h(t(b, q), q)]} > 0.$$

A.4 Proof of Lemma 3

For a feasible partial-disclosure solution that uses both actions, $b < 1$, so $c - b\Delta_B > 0$. The incentive condition can be written

$$F(t)(b\Delta_A - c) - q[1 - F(t)](c - b\Delta_B) \geq 0.$$

For $q > 0$, effort after h therefore requires $b\Delta_A - c > 0$.

The principal's cutoff condition implies

$$g_A + qg_B - (\kappa + q)t = (b\Delta_A - c) + q(c - b\Delta_B) > 0. \quad (\text{A.2})$$

Thus $t < t^S(q)$ for every feasible cutoff with $q > 0$ that uses both actions.

When $q = 0$, the incentive condition is $b\Delta_A \geq c$. The cutoff condition gives $t = (1 - b)\Delta_A/\kappa$, so the largest feasible cutoff is obtained at $b = c/\Delta_A$ and equals g_A/κ . If F has bounded support and $F(t) = 1$, the principal's expected payoff is constant as the cutoff moves further above the support. Together with (12), this establishes the stated weak monotonicity throughout the feasible set.

A.5 Proof of Proposition 2

Write

$$b(t, q) = 1 - \frac{(\kappa + q)t}{\Delta_A - q\Delta_B}.$$

Consider the set of (t, q) satisfying $q \in [0, 1]$, $t \geq 0$, $b(t, q) \in [0, 1]$, and the unnormalized incentive condition

$$F(t)[b(t, q)\Delta_A - c] + q[1 - F(t)][b(t, q)\Delta_B - c] \geq 0.$$

The cutoff condition gives

$$t \leq \frac{\Delta_A - q\Delta_B}{\kappa + q} \leq \frac{\Delta_A}{\kappa}.$$

Since F is continuous, this feasible set is closed and bounded, hence compact. Its projection $\mathcal{Q}$ is therefore compact, and for each $q \in \mathcal{Q}$ the corresponding section is compact, so the largest feasible cutoff $\tau(q)$ exists. The objective $\mathcal{W}$ is continuous because F has finite mean.

For a fixed $q \in \mathcal{Q}$, Lemma 3 shows that $\mathcal{W}$ is weakly increasing throughout the feasible section, and strictly increasing while both actions are chosen. Hence the principal selects $\tau(q)$. If F has bounded support, an all- A outcome with effort is included in this set whenever $F(t) = 1$. At a nondegenerate selected cutoff, the incentive condition must bind. Otherwise the bonus could

be reduced slightly, raising the cutoff while preserving the incentive condition and increasing the principal's payoff. Binding gives (8). For $q > 0$, combining it with (7) gives (10). The case $q = 0$ is treated directly by (13).

The reduced problem can therefore be represented by choosing $q \in \mathcal{Q}$ and then the largest feasible cutoff. Continuity on the compact feasible set gives a maximum. Comparing it with the finite payoffs in (14) and (15) proves existence for the global problem.

A.6 Proof of Proposition 3

Conditional on a committed rule specifying the action for each Z , disclosure can affect the principal's payoff only through x_A and x_B . The payoff in (18) follows from total surplus and the binding participation constraint. It is weakly increasing in x_A , strictly so when $\lambda > 0$, and linear in x_B . This gives (19).

Under full disclosure, choose

$$\frac{c}{\Delta_A} \leq b < \frac{c}{\Delta_B}$$

to implement $(x_A, x_B) = (1, 0)$, or choose $b \geq c/\Delta_B$ to implement $(1, 1)$. In the latter case the bonus exceeds one, which is feasible because the managerial action has already been committed and the bonus cannot distort action choice. The fixed payment adjusts to satisfy the participation constraint. Thus full disclosure implements the payoff-maximizing effort choice following each action.

A.7 Proof of Proposition 4

Take any outcome attained with contractible signals. For each action a , let x_a and T_a denote the probability that the agent exerts effort and the expected transfer defined in Section 6.2. Under full action disclosure, offer the contract in (22). Since $b_a \Delta_a = c$, the agent is indifferent between effort and no effort. Exerting effort with probability x_a is therefore a best response under the convention stated in Section 3. Equivalently, an independent public lottery can implement this probability while the action remains fully revealed. Expected compensation is

$$s_a + b_a(p^0 + x_a \Delta_a) = T_a.$$

Thus the probability that the agent exerts effort, expected transfer, and agent utility conditional on each action are preserved. The payoffs in (20) and (21) show that the principal's payoff from each action is preserved for every Z . Her action strategy and the entire outcome are therefore unchanged.

A.8 Proof of Proposition 5

For the two-signal statement, group all original signals according to whether the agent exerts effort. The proof of Lemma 1 applies with γ_A and γ_B in place of $b\Delta_A$ and $b\Delta_B$. The probabilities that the agent exerts effort conditional on each action and the principal's payoff from each action are preserved.

For the one-sided statement, let

$$\alpha = \Pr(h \mid A) > 0, \quad \beta = \Pr(h \mid B) = \alpha q, \quad q \in [0, 1].$$

The principal's payoff relevant for action choice is $\alpha\chi_A$ under A and

$$\kappa Z + \alpha q(\chi_B + Z)$$

under B . She therefore chooses A if and only if $Z \leq t(\alpha, q)$, where

$$t(\alpha, q) = \frac{\alpha(\chi_A - q\chi_B)}{\kappa + \alpha q}.$$

The assumed ordering gives

$$\frac{\partial t}{\partial \alpha} = \frac{\kappa(\chi_A - q\chi_B)}{(\kappa + \alpha q)^2} > 0.$$

Let $d_a = \gamma_a + \chi_a$ be the expected increase in output generated by effort under action a . Using $M(t)$ defined in Section 5.2, define

$$H_d(t) = \int_t^\infty (z + d_B - c) dF(z).$$

Apart from a constant independent of α , the principal's expected payoff after using the fixed payment to satisfy the participation constraint is

$$V(\alpha, q) = \kappa M(t) + \alpha[(d_A - c)F(t) + qH_d(t)].$$

Differentiation and the cutoff condition give

$$\begin{aligned} V_\alpha &= (d_A - c)F(t) + qH_d(t) \\ &\quad + \alpha f(t)t_\alpha[(\gamma_A - c) + q(c - \gamma_B)]. \end{aligned}$$

The second line is nonnegative. The first line can be written

$$\begin{aligned} &(d_A - c)F(t) + qH_d(t) \\ &= F(t)(\gamma_A - c) + q[1 - F(t)](\gamma_B - c) \\ &\quad + F(t)\chi_A + q[1 - F(t)]\chi_B + qM(t). \end{aligned}$$

The first two terms on the right are nonnegative because effort after h is optimal. The remaining terms are strictly positive in a nondegenerate equilibrium because $\chi_A > \chi_B \geq 0$. Hence $V_\alpha > 0$.

As α rises with q fixed, the cutoff rises and h becomes weakly more indicative of A , so effort after h remains optimal. At $\alpha = 1$, ℓ is generated only by B , and no effort after ℓ is optimal because $\gamma_B < c$. The endpoint is feasible and strictly improves the principal's payoff unless the policy already

has $\alpha = 1$.

A.9 Proof of Lemma 4

Under an arbitrary disclosure policy, let $\rho_j(\sigma)$ be the probability that the agent chooses effort e_j after signal σ . For each managerial action a , define

$$q_{aj} = \int \rho_j(\sigma) \pi(d\sigma | a).$$

Construct a policy with signals $1, \dots, K$ satisfying $\Pr(j | a) = q_{aj}$.

Let P be the unconditional distribution of the original signals, let $\mu(\sigma)$ be the posterior probability of A , and write

$$D(\sigma) = \mu(\sigma)\Delta_A + [1 - \mu(\sigma)]\Delta_B.$$

Whenever $\rho_j(\sigma) > 0$, optimality of e_j implies, for every alternative e_k ,

$$b(e_j - e_k)D(\sigma) - [C(e_j) - C(e_k)] \geq 0.$$

Multiplying by $\rho_j(\sigma)$ and integrating with respect to P gives

$$\Pr(j)\{b(e_j - e_k)D_j - [C(e_j) - C(e_k)]\} \geq 0,$$

where D_j is the expected increase in the probability of $y = 1$ from a one-unit increase in the agent's effort after the constructed signal j . Thus e_j is optimal after that signal whenever it occurs with positive probability.

The construction preserves the probability of every effort level conditional on each managerial action. Expected output, compensation, effort cost, and the principal's payoff from each action for every Z are therefore preserved. The principal's action choice and both parties' ex ante payoffs are unchanged.

A.10 Proof of Proposition 6

Suppose the cutoff is t . The disclosure policy and Bayes' rule determine $\bar{e}_A(t)$ and $\bar{e}_B(t)$. Excluding terms common across actions, the principal's expected payoff from A is

$$(1 - b)\Delta_A \bar{e}_A(t).$$

Her expected payoff from B is

$$(1 - b)\Delta_B \bar{e}_B(t) + Z[\kappa + \bar{e}_B(t)].$$

The payoff difference is strictly decreasing in Z . The principal chooses A if and only if

$$Z \leq \frac{(1-b)[\Delta_A \bar{e}_A(t) - \Delta_B \bar{e}_B(t)]}{\kappa + \bar{e}_B(t)}.$$

The cutoff is consistent exactly when it equals the right side, which gives (23).

A.11 Proof of Proposition 7

Using (24),

$$\begin{aligned} \mathcal{V}(\mathbf{q}) - I(\mathbf{q}) &= p^0 - \bar{u} + \sum_{i=1}^n \int_{Z_i(\mathbf{q})} [\phi_i(z) + q_i R_i(z)] dF(z) \\ &= p^0 - \bar{u} + \int \max_i \{\phi_i(z) + q_i R_i(z)\} dF(z). \end{aligned}$$

At a differentiability point where the maximizing action is unique almost surely, the envelope theorem gives

$$\frac{\partial}{\partial q_k} [\mathcal{V}(\mathbf{q}) - I(\mathbf{q})] = \int_{Z_k(\mathbf{q})} R_k(z) dF(z).$$

Rearranging proves (25).

A.12 Proof of Proposition 8

At a nondegenerate optimum, the effort incentive condition binds. The agent's expected utility is therefore $s + bp^0$. Since the participation constraint binds in the unrestricted problem, $s = \bar{u} - bp^0$.

If $p^0 = 0$, the unrestricted partial-disclosure optimum uses $s = \bar{u} \geq 0$ and is feasible under limited liability. Since the limited-liability feasible set is a subset of the unrestricted set, the same outcome remains optimal.

For the second part, observe that the effort incentive condition and the principal's cutoff condition do not depend on p^0 . The feasible design set is therefore unchanged as p^0 varies. The compactness argument in the proof of Proposition 2 implies that a strict global optimum has a positive payoff gap outside any sufficiently small neighbourhood. Under limited liability, both the interior objective and the candidate all- B payoffs vary continuously with p^0 . The maximum theorem therefore preserves this gap for all sufficiently small $p^0 > 0$. With $\bar{u} = 0$, limited liability gives $s = 0$, and the agent's rent is $bp^0 > 0$.

Declaration of generative AI and AI-assisted technologies in the manuscript preparation process

During the preparation of this work, the author used OpenAI’s ChatGPT and Anthropic’s Claude to assist with language editing, exposition, and the review and organization of related literature. The author reviewed and edited all AI-assisted output and takes full responsibility for the content of the article.

References

- Angwin, Duncan N., Kamel Mellahi, Emanuel Gomes, and Emmanuel Peter. How communication approaches impact mergers and acquisitions outcomes. *The International Journal of Human Resource Management*, 27(20):2370–2397, 2016.
- Babichenko, Yakov, Inbal Talgam-Cohen, Haifeng Xu, and Konstantin Zabarnyi. Information design in the principal-agent problem. *Games and Economic Behavior*, 155:55–69, 2026.
- Baliga, Sandeep, Luis C. Corchón, and Tomas Sjöström. The theory of implementation when the planner is a player. *Journal of Economic Theory*, 77(1):15–33, 1997.
- Baliga, Sandeep, and Tomas Sjöström. Interactive implementation. *Games and Economic Behavior*, 27(1):38–63, 1999.
- Bergemann, Dirk, and Stephen Morris. Bayes correlated equilibrium and the comparison of information structures in games. *Theoretical Economics*, 11(2):487–522, 2016.
- Blanes i Vidal, Jordi, and Marc Möller. When should leaders share information with their subordinates? *Journal of Economics & Management Strategy*, 16(2):251–283, 2007.
- Boleslavsky, Raphael, and Kyungmin Kim. Bayesian persuasion and moral hazard. Working paper, 2021.
- Catonini, Emiliano, and Sergey Stepanov. On the optimality of full disclosure. *RAND Journal of Economics*, 57(2):300–314, 2026.
- Chen, Deqiu, Huasheng Gao, and Yujing Ma. Human capital-driven acquisition: Evidence from the inevitable disclosure doctrine. *Management Science*, 67(8):4643–4664, 2021.
- Curello, Gregorio, Qianjun Lyu, and Yimeng Zhang. Complete contracts under incomplete information. Working paper, May 2026.
- Holmström, Bengt. Moral hazard and observability. *Bell Journal of Economics*, 10(1):74–91, 1979.
- Jehiel, Philippe. On transparency in organizations. *Review of Economic Studies*, 82(2):736–761, 2015.

- Jost, Peter-Jürgen. On the role of commitment in a principal-agent relationship with an informed principal. *Journal of Economic Theory*, 68(2):510–530, 1996.
- Kamenica, Emir, and Matthew Gentzkow. Bayesian persuasion. *American Economic Review*, 101(6):2590–2615, 2011.
- Kwak, Changhyun. Optimal information disclosure with moral hazard. Working paper, 2022.
- Lewis, Tracy R., and David E. M. Sappington. Information management in incentive problems. *Journal of Political Economy*, 105(4):796–821, 1997.
- Li, Anqi, and Ming Yang. Optimal incentive contract with endogenous monitoring technology. *Theoretical Economics*, 15(3):1135–1173, 2020.
- Myerson, Roger B. Optimal coordination mechanisms in generalized principal-agent problems. *Journal of Mathematical Economics*, 10(1):67–81, 1982.
- Schweiger, David M., and Angelo S. DeNisi. Communication with employees following a merger: A longitudinal field experiment. *Academy of Management Journal*, 34(1):110–135, 1991.
- Sobel, Joel. Information control in the principal-agent problem. *International Economic Review*, 34(2):259–269, 1993.
- Sonenshein, Scott, and Utpal Dholakia. Explaining employee engagement with strategic change implementation: A meaning-making approach. *Organization Science*, 23(1):1–23, 2012.
- Tan, Teck Yong. Optimal transparency of monitoring capability. *Journal of Economic Theory*, 209:105620, 2023.
- Younge, Kenneth A., Tony W. Tong, and Lee Fleming. How anticipated employee mobility affects acquisition likelihood: Evidence from a natural experiment. *Strategic Management Journal*, 36(5):686–708, 2015.