

The Value of Non-Value-Maximizing Managers

Finance Working Paper N° 1143/2026

June 2026

Pierre Chaigneau

Queen's University and ECGI

Alex Edmans

London Business School, CEPR and ECGI

Nicolas Sahuguet

HEC Montréal

© Pierre Chaigneau, Alex Edmans and Nicolas Sahuguet 2026. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

This paper can be downloaded without charge from:
http://ssrn.com/abstract_id=6895618

www.ecgi.global/content/working-papers

ECGI Working Paper Series in Finance

The Value of Non-Value-Maximizing Managers

Working Paper N° 1143/2026

June 2026

Pierre Chaigneau
Alex Edmans
Nicolas Sahuguet

Abstract

Companies pursue multiple goals, such as financial returns and societal impact. Conventional wisdom suggests they should appoint managers who share investors' preferences across these goals. We show that, when performance measurement is asymmetric, appointing a manager who is biased towards the worse-measured activity ("impact") may instead be optimal. Doing so allows for stronger incentives on the better-measured activity ("profits") without distorting resource allocation. Strong financial incentives need not indicate that the manager lacks intrinsic motivation; instead, they counteract his intrinsic motivation for impact. Similarly, strong financial incentives need not imply that a mission-focused firm is greenwashing; it pursues its mission through a biased manager rather than impact incentives. In a market equilibrium with a continuum of managers, we characterize the optimal level of bias and how it affects the level and structure of pay. The framework offers new insights into executive incentives, corporate purpose, and organizational culture.

Keywords: executive compensation, incentives, sustainability, ESG, private benefits, managerial preferences

Pierre Chaigneau

Associate Professor of Finance & Commerce '77 Fellow of Finance
Queen's University
143 Union Street West
Kingston, ON K7L 2P3, Canada
e-mail: pierre.chaigneau@queensu.ca

Alex Edmans*

Professor of Finance
London Business School
Regent's Park
London, NW1 4SA, United Kingdom
e-mail: aedmans@london.edu

Nicolas Sahuguet

Professor
HEC Montréal
3000, chemin de la Côte-Sainte-Catherine
Montréal (Québec), Canada
e-mail: nicolas.sahuguet@hec.ca

*Corresponding Author

The Value of Non-Value-Maximizing Managers

Pierre Chaigneau*

Queen’s and ECGI

Alex Edmans[†]

LBS, CEPR, and ECGI

Nicolas Sahuguet[‡]

HEC Montréal and CEPR

June 7, 2026

Abstract

Companies pursue multiple goals, such as financial returns and societal impact. Conventional wisdom suggests they should appoint managers who share investors’ preferences across these goals. We show that, when performance measurement is asymmetric, appointing a manager who is biased towards the worse-measured activity (“impact”) may instead be optimal. Doing so allows for stronger incentives on the better-measured activity (“profits”) without distorting resource allocation. Strong financial incentives need not indicate that the manager lacks intrinsic motivation; instead, they counteract his intrinsic motivation for impact. Similarly, strong financial incentives need not imply that a mission-focused firm is greenwashing; it pursues its mission through a biased manager rather than impact incentives. In a market equilibrium with a continuum of managers, we characterize the optimal level of bias and how it affects the level and structure of pay. The framework offers new insights into executive incentives, corporate purpose, and organizational culture.

Keywords: executive compensation, incentives, sustainability, ESG, private benefits, managerial preferences.

*Address: Smith School of Business, Goodes Hall, 143 Union Street West, K7L 2P3, Kingston, ON, Canada. Email: pierre.chaigneau@queensu.ca.

[†]Address: Regent’s Park, London NW1 4SA, United Kingdom. Email: aedmans@london.edu.

[‡]Address: Applied Economics Department, HEC Montréal, 3000 chemin de la côte Sainte Catherine, H3T 2A7 Montréal, QC, Canada. Email: nicolas.sahuguet@hec.ca.

1 Introduction

Organizations routinely face decisions that require trading off multiple dimensions of performance. A board may care about both financial returns and environmental and social (“ES”) impact, a hospital about treatment and preventive care, and a university about research output and teaching quality. These decisions may involve allocating resources such as money or time, setting the organization’s priorities, or determining its strategy.

It might seem optimal to appoint an unbiased leader who shares the organization’s overall objective and can be relied upon to make these decisions in pursuit of that objective. However, some companies select CEOs who appear to be biased towards one particular goal. Several executives, such as BlackRock’s Larry Fink and former Disney boss Bob Chapek, have been criticized for prioritizing ES impact over financial returns – for example through efforts to combat climate change, promote diversity, equity, and inclusion (“DEI”), or support social justice initiatives. Such actions, and even the appointment of these CEOs, are sometimes dismissed as “woke capitalism.” Other leaders, such as Tesla’s Elon Musk or Meta’s Mark Zuckerberg, are viewed as prioritizing technological innovation and ambitious scientific vision even at the expense of shareholder value.

This paper shows that such practices may be optimal. Central to our result is the observation that, across different dimensions, measured performance may diverge from actual output to varying degrees. For example, shareholder value can be proxied by the stock price in an efficient market, but many aspects of ES impact are qualitative and third-party quantitative ratings are inconsistent (Berg, Kölbel, and Rigobon, 2022). Focusing contractual incentives on the better-measured dimensions may lead the CEO to neglect the worse-measured ones (Holmström and Milgrom, 1991). Hiring managers whose intrinsic preferences are biased towards the latter allows the board to use high-powered incentives on the former, while relying on the manager’s bias to prevent excessive neglect of the latter. Thus, appointing agents whose preferences differ from the principal’s objective can help achieve her objective: non-value-maximizing managers can be value-maximizing.

We build a parsimonious principal-agent model in which a manager allocates a fixed resource across two activities. For clarity, we will refer to the better-measured activity as “profits” and the worse-measured activity as “impact”, although the model admits many other interpretations. Even under this interpretation, “impact” has several applications, such as ES impact, scientific innovation and technological progress, or brand strength and customer reputation, and the model also allows “impact” to be better measured than “profits.” The manager privately observes the productivity of each activity. Output of an activity is increasing in both productivity and the resources allocated to it. While output is non-contractible, an imperfect measure of performance is available and contractible for each activity. The principal’s overall objective is the total output

from both activities, minus the manager’s pay and an adjustment cost of deviating from a default allocation. Like output, this objective is also non-contractible. The manager maximizes his pay minus the adjustment cost of deviating from his default allocation; the latter reflects his preferences and may differ from the principal’s default allocation.

Strong incentives on a given dimension induce the manager to respond to his private information about that dimension’s productivity, rather than adhere to his default allocation. However, if measured performance diverges from true output, strong incentives also lead the manager to prioritize the former at the expense of the latter (Kerr (1975), Baker (1992)): to allocate resources towards an activity when doing so increases measured performance, even if productivity and thus the contribution to actual output is low. Because profits are more precisely measured, optimal incentive contracts place greater weight on profits and less weight on impact. This asymmetry, however, leads the manager to allocate too many resources to profits from the principal’s perspective. Reducing incentives for profits is not a solution, as it would reduce the manager’s responsiveness to his private information. Increasing incentives for impact is not a solution either, as it would lead him to improve measured impact rather than actual impact – to “hit the target, miss the point.”

We show how this distortion can be mitigated through managerial selection. When impact is poorly measured, the principal optimally appoints a manager who values it more than she does. This bias allows the principal to offer high-powered incentives for profits, where they are most effective, without causing the manager to neglect impact. Note that bias is different from intrinsic motivation. In a single-activity model, it is well-known that an intrinsically motivated manager (e.g. one with a private benefit from effort or output) is valuable because it attenuates the agency problem (Akerlof and Kranton (2000, 2005), Besley and Ghatak (2005), Carlin and Gervais (2009)). In our setting, there are two activities that conflict with each other. Thus, while bias towards impact creates “intrinsic motivation” for impact, it is at the expense of profits. Indeed, in a first-best world where the principal can observe productivity and the resource allocation decision, bias is unambiguously costly. In addition, unlike in single-activity models where the activity (e.g. effort) is always beneficial to the principal, and so intrinsic motivation is always desirable, here the manager may pursue excessive impact.

Our results also contrast the standard multitasking model which predicts low-powered incentives across all tasks when tasks are substitutes (Holmström and Milgrom (1991)). In a multitasking model with imperfect performance measurement, Prendergast (2008) shows that the principal can hire “biased agents” whose bias is for a particular task. The implication is to use job design by allocating different types of agents to different tasks. In our model, the agent only has one task – allocating resources – so that job design is not a solution; it is thus particularly applicable to CEOs, who are responsible for all of a company’s major activities. Instead, the solution is a combination of bias and financial incentives.

The model provides implications for the link between bias and incentives. We find that increasing the manager’s bias towards impact not only reduces the optimal incentives for impact, but also increases his incentives for profits. This result rationalizes some otherwise puzzling practices. Many companies state a purpose beyond shareholder value, yet the vast majority of executive incentives are tied to the stock price (Bebchuk and Tallarita, 2023), leading to criticism that purpose statements are “greenwashing”. Vladasel et al. (2024) find that many social enterprises simultaneously offer a social mission and monetary rewards for (more easily measured) commercial tasks. However, given the difficulties in measuring impact, the optimal way to pursue it may be to appoint an impact-conscious manager – in which case strong financial incentives are needed to ensure that resources are not excessively allocated to impact at the expense of profitability. Separately, strong financial incentives are often criticized on the grounds that intrinsic motivation should be sufficient to motivate managerial effort: if the CEO requires incentives to deliver exceptional performance, the board has hired the wrong CEO (see the survey of Edmans, Gosling, and Jenter (2023)). Our model shows that financial incentives are fully consistent with an intrinsically motivated CEO. Indeed, the CEO may be intrinsically motivated towards an action (impact) that creates value and enhances the principal’s objective up to a point. Financial incentives are needed to ensure that he does not pursue impact beyond that point.

Our framework also yields predictions for the types of companies that should appoint biased managers. Bias is more valuable when performance measurement is particularly asymmetric, because financial incentives will be more distortionary. It is also greater when the manager’s private information is more important, since financial incentives will be stronger and there is more asymmetry to correct. Bias is less valuable when the manager faces a high adjustment cost from deviating from his default allocation (as he is less responsive to incentives and so incentives are less distortionary), and when the firm faces a high adjustment cost from its default allocation (as a biased manager with a different default allocation is less desirable).

We next extend the analysis to an economy with heterogeneous firms and managers. Firms differ in the quality of their performance measurement, while managers differ in their biases. In a matching equilibrium, we show that firms with worse measurement of impact hire managers who are more biased toward impact. Up to a point, managers who are more biased towards impact earn greater rents: this bias is valuable when firms measure impact poorly and so incentives for impact are less effective. Managers who are intrinsically biased away from shareholders’ objective may have greater potential to achieve shareholders’ objective, given the endogenous response of financial incentives, and thus earn higher pay. However, when the bias is too high, the manager destroys value as his chosen allocation deviates too much from the principal’s. Unlike in the single-firm model, a manager who is more biased towards profits can actually receive stronger incentives for profits. This is because he is matched with a firm that measures impact well

and thus offers high impact incentives. The firm thus offers high profit incentives to reduce the distortionary effect of high impact incentives.

The manager's preference for a particular activity may stem from a number of sources. One is personal preferences, which may be shaped by his background and upbringing. A second is organizational factors such as corporate culture or social norms. For example, if a firm has issued a purpose statement committing to impact, has a historical reputation for delivering it, or employs colleagues who value it, the manager may face social or reputational pressure to deliver impact. Interestingly, our model predicts that improved measurement of impact (e.g. through superior ES reporting standards) may optimally shift corporate culture away from impact, as impact can instead be incentivized through compensation contracts. Conversely, improvements in financial reporting, and thus the effectiveness of profit incentives, may lead boards to place greater emphasis on impact in public statements, to encourage managers to deliver it. The model thus rationalizes apparent overcommitment to impact (e.g. through mission statements), even if impact is poorly measured and may be detrimental to shareholder value. Such a commitment may reflect an optimal response to the superior measurement of profits.

This paper principally contributes to two literatures: one on performance measurement and a second on intrinsically motivated agents. Starting with the former, Holmström and Milgrom (1991) show that when some dimensions of performance are easier to measure than others, incentive pay distorts the allocation of effort toward the better-measured dimensions. They propose job design as a solution: assigning easy-to-measure tasks to one agent, who is given strong incentives for all tasks, and difficult-to-measure tasks to another agent alongside weak incentives. Baker (1992) considers a single performance measure and finds that optimal incentives are weaker when this measure is less closely aligned with the principal's objective. Moving to the latter, Akerlof and Kranton (2000, 2005) and Fischer and Huddart (2008) show that an agent who identifies with the principal's goals or follows norms that promote high effort does not need strong financial incentives. Besley and Ghatak (2005, 2006, 2018) show that, in a market equilibrium, firms hire agents whose preferences accord with their mission, allowing the agency problem to be mitigated at lower cost. We instead show that asymmetric measurement can make congruence suboptimal: the principal may prefer a manager who is biased toward the worse-measured activity.

Our paper is closest to work showing that some bias can be valuable. Prendergast (2007) shows that appointing biased agents (those who place greater weight on a particular dimension than the principal) can increase effort in a model without financial incentives, and so firms should choose agents whose biases are suited to the task at hand. Prendergast (2008) introduces financial incentives on the single aggregate performance measure, and finds that incentives are lower for more biased (intrinsically motivated) agents. The firm hires a different agent to focus on each task, with a bias suited to that task. We consider a single resource allocation decision across

two activities. Each activity is separately measured, with performance measures that differ in their accuracy, and thus separately incentivized. This generates asymmetric incentives across activities and, in turn, distorts resource allocation toward the better-measured dimension. Bias therefore plays a different role: it counteracts the incentive-induced misallocation of resources, thereby allowing stronger incentives, rather than substituting for weak incentives that arise because performance imperfectly captures output. In Prendergast (2008), bias reduces the need for incentives; in our setting, it improves their effectiveness.

Bénabou and Tirole (2016) study a setting in which one task is contractible and the other is not, and must therefore be induced through intrinsic motivation. When intrinsic motivation differs across agents and is privately observed, they show that agents with greater motivation for “impact” than other agents select weakly lower incentives for “profits” (using our terminology for simplicity). Because effort across tasks is substitutable, stronger profit incentives draw effort away from the intrinsically valued impact task; more impact-motivated agents therefore value such incentives less at the margin and instead prefer fixed pay. In contrast, we show that impact-motivated managers receive stronger profit incentives, because their impact motivation mitigates distortions induced by profit incentives.

Finally, our paper is related to research on the consequences of managers focusing on narrow financial metrics. Aghion and Stein (2008) and Ben-David and Chinco (2026) study settings in which managerial decisions are shaped by metrics such as profit margins or short-term EPS, even though they are imperfect proxies for firm value. We show that tying incentive pay to these proxies may be efficient, despite their incompleteness, if managers have intrinsic preferences towards other performance dimensions.

2 Model

We consider a principal-agent model of resource allocation. A principal hires an agent (manager) who receives private information about the productivity of potential allocations of the firm’s resources.¹ The contracting principal may be a board acting on behalf of shareholders, an investor, or the firm itself; for ease of exposition, we often refer to the principal as “the firm”. The firm produces unverifiable output V_i on two dimensions $i = 1, 2$. On each dimension, performance is measured by the verifiable metric P_i . Output and performance are affected by a non-contractible allocation decision, a . This decision can reflect the allocation of time, money or other resources, or the choice of strategy. As in Baker (1992), output and performance are

¹Private information may be interpreted literally; alternatively, it may result from the manager’s expertise, thus allowing him to use public information to make superior resource allocation decisions. For example, an experienced CEO may observe industry trends and use them to determine the optimal scale of investment.

also affected by shocks $\epsilon \equiv \{\epsilon_1^v, \epsilon_2^v, \epsilon_1^p, \epsilon_2^p\}$ which are privately observed by the manager:

$$V_1(a, \epsilon) = a\epsilon_1^v \quad \text{and} \quad V_2(a, \epsilon) = (1 - a)\epsilon_2^v \quad (1)$$

$$P_1(a, \epsilon) = a\epsilon_1^p \quad \text{and} \quad P_2(a, \epsilon) = (1 - a)\epsilon_2^p \quad (2)$$

For $i \in \{1, 2\}$ and $j \in \{v, p\}$, the random variable ϵ_i^j has a standard deviation σ_i^j , and a mean normalized to $\mathbb{E}[\epsilon_i^j] = 1$. The correlation coefficient between output and measured performance on dimension i is: $\rho_i \equiv \frac{\text{cov}(\epsilon_i^v, \epsilon_i^p)}{\sigma_i^v \sigma_i^p}$. Shocks to performance and output are independent across dimensions, so that:²

$$\text{cov}(\epsilon_i^p, \epsilon_j^v) = 0 \quad \text{and} \quad \text{cov}(\epsilon_i^p, \epsilon_j^p) = 0 \quad \forall i \neq j.$$

The manager thus makes a single resource allocation decision, a , which affects both outputs and performance measures.³ This differs from the multitasking framework, in which a manager with multiple tasks chooses separate effort levels for each task. While both frameworks allow for allocations across tasks to be substitutes, they differ in several important ways. First, in a multitasking model, individual tasks can in principle be allocated to different agents (Holmström and Milgrom (1991); Prendergast (2007, 2008)). By contrast, the resource allocation model involves a single task, and thus is most applicable to CEOs who are responsible for all major activities. Second, multitasking models typically assume that the principal knows which actions are desirable (higher effort on each task) so inefficiencies arise from the principal's inability to induce them. In such settings, intrinsic motivation can help align the agent's behavior with the principal's objectives by increasing effort on desired tasks. Here, the principal does not know which allocation is optimal, as it depends on the agent's private information. As a result, bias plays a fundamentally different role: rather than helping to implement the principal's preferred actions, it can distort the agent's decision rule, leading to inefficient choices even in the absence of incentive frictions. The desirability of bias is thus a priori unclear. Third, with a single action, preferences over allocations can be summarized by a single parameter, making it straightforward to compare the principal's and manager's preferred allocations.

The firm's objective function is:

$$\mathbb{E} \left[V_1(a, \epsilon) + V_2(a, \epsilon) - W(a, \epsilon) - \frac{k}{2} (a - \alpha_P)^2 \right] \quad (3)$$

where $W(a, \epsilon)$ is the manager's payment, and $k > 0$. The last term captures the cost to the

²This does not imply that performance will be independent across dimensions: increasing a improves V_1 at the expense of V_2 . Instead, we are only assuming that shocks to performance are independent across dimensions.

³For tractability, we do not restrict $a \in [0, 1]$. Note that $a < 0$ could be interpreted as the firm not spending any resources on impact, but instead having a harmful impact to increase profitability, whereas $a > 1$ could be interpreted as the firm raising financing to spend more than its existing resources on impact.

principal of deviating from a baseline resource allocation α_P , which may reflect adjustment costs (e.g., if α_P represents a historical or approved allocation), or board or investor preferences.

The firm's objective is noncontractible, since outputs are noncontractible. For example, the outputs may be profit and impact, and shareholders care about both, as in models of responsible investing. Alternatively, the objective function is long-term shareholder value and V_1 and V_2 are dimensions that both contribute to shareholder value. If long-term shareholder value is not reflected in the stock price for many years (e.g. one dimension is investment in a new technology that the market is unable to value in the short term), and there is a limit to how long the firm can escrow the manager's equity (e.g. due to his consumption needs), then shareholder value is not contractible. If the principal's objective were contractible, the agency problem would be simple to solve as the principal would contract on it.

As in Baker (1992), the firm offers a linear incentive contract $\{S, b_1, b_2\}$ such that the manager's payment is equal to:

$$W(a, \epsilon) = S + b_1 P_1(a, \epsilon) + b_2 P_2(a, \epsilon) \quad (4)$$

Linear contracts are a standard assumption in the literatures on performance measurement (Baker (1992)) and intrinsic motivation (Fischer and Huddart (2008), Prendergast (2008), Bénabou and Tirole (2016)). They allow incentives for a task to be captured by a single parameter, b_i , allowing for clear results on the link between incentives and bias (the focus of this paper), as well as straightforward comparative static analysis. The Online Appendix provides one potential microfoundation: if the manager can induce mean-preserving spreads or contractions of the performance measure distributions, then he can game any nonlinear contract. Only linear contracts are robust to such manipulation, as also shown by Barron, Georgiadis, and Swinkels (2020).

The manager's reservation utility is given by $\bar{U}$. Given a contract $\{S, b_1, b_2\}$, the manager's information ϵ , and resource allocation a , his expected utility is:

$$\mathbb{E} \left[S + b_1 P_1(a, \epsilon) + b_2 P_2(a, \epsilon) - \frac{c}{2} (a - \alpha_M)^2 \right], \quad (5)$$

where $c > 0$ parameterizes the cost to the manager of deviating from his preferred resource allocation α_M . The manager has preferences over the resource allocation decision, similar to models of motivated agents (e.g., Akerlof and Kranton (2005); Besley and Ghatak (2005)). Empirically, there is a large body of evidence that agents have nonpecuniary preferences over actions, and that such preferences are heterogeneous across individuals (e.g., Wrzesniewski et al. (1997); Cassar and Meier (2018), Fehr and Charness (2025)). There is a continuum of different managers that the firm can hire, each of whom has a different α_M , where $\alpha_M \in [\underline{\alpha}_M, \overline{\alpha}_M]$. If the firm and manager's baseline resource allocations are the same ($\alpha_M = \alpha_P$), we say that the

manager is “congruent” or “unbiased”. We use “bias” to refer to the difference between α_M and α_P , i.e. if $\alpha_M > \alpha_P$, the manager is biased towards dimension 1.

The firm chooses the contract $\{S, b_1, b_2\}$ to maximize its objective function in equation (3), subject to the manager’s incentive constraint (he chooses a to maximize (5)) and participation constraint (his expected utility is at least $\bar{U}$).

3 Analysis

This section solves for both the type of manager that the firm hires and the compensation contract that it offers. We begin by characterizing the first-best, which arises when the resource allocation a is contractible and there is no asymmetric information (i.e. the principal observes the state ϵ). Then, the principal can dictate a state-contingent allocation $a(\epsilon)$ and pays a fixed wage w . She chooses the resource allocation $a(\epsilon)$ and the manager’s type α_M to maximize her expected payoff subject to the manager’s participation constraint:

$$\max_{a(\epsilon), \alpha_M} \mathbb{E} \left[a \epsilon_1^v + (1 - a) \epsilon_2^v - w - \frac{k}{2} (a - \alpha_P)^2 \right] \quad (6)$$

$$\text{s.t.} \quad \mathbb{E} \left[w - \frac{c}{2} (a - \alpha_M)^2 \right] \geq \bar{U}. \quad (7)$$

The participation constraint will bind at the optimum. Substituting for w from the binding participation constraint into the objective function, and optimizing pointwise over a for each realization of ϵ , yields the first-best state-contingent allocation:

$$a^{FB}(\epsilon) = \frac{c\alpha_M + k\alpha_P + \epsilon_1^v - \epsilon_2^v}{c + k}. \quad (8)$$

Substituting back into (6) and optimizing over α_M shows that the principal hires a congruent manager, i.e. $\alpha_M = \alpha_P$. Intuitively, this minimizes the manager’s expected disutility from deviation, which must be compensated to satisfy his participation constraint; there is no benefit to hiring a biased manager as the allocation a can be dictated anyway. Then, the first-best allocation is:

$$a^{FB}(\epsilon) = \alpha_P + \frac{\epsilon_1^v - \epsilon_2^v}{c + k}. \quad (9)$$

The allocation depends on the output shocks ϵ_1^v and ϵ_2^v , not the performance shocks ϵ_1^p and ϵ_2^p . The wedge between output and measured performance, which is the main friction in the second-best, is irrelevant in the first-best when the principal observes ϵ and directly controls a .

We now study the equilibrium outcome when the resource allocation decision is non-contractible, i.e. the second-best outcome. We solve the model backwards, starting with the manager’s resource allocation decision, then deriving the optimal contract, and finally solving for the man-

ager's bias α_M that the firm prefers. We start by deriving several preliminary results.

Lemma 1 *For a given contract $\{S, b_1, b_2\}$, the manager chooses the following resource allocation:*

$$a = \alpha_M + \frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c}.$$

The optimally chosen resource allocation is additive in two terms. The first is the manager's baseline resource allocation α_M , which he would choose in the absence of financial incentives. The second term captures the deviation from this baseline level due to his information and incentives, scaled by the deviation cost c . The higher the sensitivity of pay to performance on dimension 1 (b_1) and the more allocating resources to this dimension improves performance (ϵ_1^p), the more the manager allocates resources to it by increasing a . Importantly, his allocation decision depends on the effect on performance ϵ_1^p rather than the effect on output ϵ_1^v . The distinction between ϵ_1^p and ϵ_1^v means that it is suboptimal to have strong incentives for a poorly-measured dimension: the manager may increase a if ϵ_1^p is high, even if ϵ_1^v is low. For example, a bank manager may increase the number of loans even if the quality of such loans is poor, a hospital may conduct more tests even if they do not improve patient health, and a professor may design a course to improve teaching evaluations at the expense of learning quality. Likewise, the higher b_2 and ϵ_2^p , the more the manager allocates resources to the second dimension by decreasing a .

Lemma 2 determines the optimal incentives for both dimensions of performance. The fixed wage S is set to keep the manager at his reservation utility given the incentives derived in Lemma 2.

Lemma 2 *The optimal sensitivity of pay to performance on both dimensions is:*

$$b_1 = \frac{c \left(\rho_1 \sigma_1^p \sigma_1^v \left(\sigma_2^{p2} + 1 \right) + \rho_2 \sigma_2^p \sigma_2^v + \sigma_2^{p2} k (\alpha_P - \alpha_M) \right)}{(c + k) \left(\sigma_1^{p2} \sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2} \right)} \quad (10)$$

$$b_2 = \frac{c \left(\rho_2 \sigma_2^p \sigma_2^v \left(\sigma_1^{p2} + 1 \right) + \rho_1 \sigma_1^p \sigma_1^v + \sigma_1^{p2} k (\alpha_M - \alpha_P) \right)}{(c + k) \left(\sigma_1^{p2} \sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2} \right)}. \quad (11)$$

Incentives b_i are increasing in measurement quality ρ_i and ρ_j , output volatility σ_i^v and σ_j^v , and the manager's deviation cost c .

Lemma 2 shows that incentives for dimension i are additive in three terms. First, it is increasing in $\rho_i \sigma_i^v$, similar to Baker (1992).⁴ Pay is more sensitive to performance on a dimension that is well-measured (high ρ_i). It is also increasing in uncertainty σ_i^v about the production

⁴While the first term contains $\rho_i \sigma_i^p \sigma_i^v$, performance variability σ_i^p enters elsewhere in b_i so we do not have clear comparative statics.

technology, as this makes it particularly valuable to make the manager responsive to his private information.

Second, incentives b_i for dimension i are increasing in $\rho_j \sigma_j^v$, as this increases incentives for dimension j . This is a standard effect when tasks are substitutes so that the balance of incentives matters. Thus, improvements in performance measurement on one dimension increase incentives for both dimensions – a spillover effect.

Third, incentives depend on the manager's bias, $\alpha_M - \alpha_P$. When the manager is biased towards dimension 1 ($\alpha_M > \alpha_P$), incentives decrease on dimension 1 and increase on dimension 2; the reverse is true if $\alpha_M < \alpha_P$. These three terms are scaled by c : when the manager faces a higher deviation cost, incentives need to be stronger to induce him to follow his private information.

Proposition 1 delivers the main result of the paper: the firm's optimal choice of the manager's preferences α_M .⁵

Proposition 1 *The firm appoints a manager with the following preferred resource allocation:*

$$\alpha_M = \max \left\{ \min \left\{ \alpha_M^*, \overline{\alpha_M} \right\}, \underline{\alpha_M} \right\} \quad \text{where} \quad \alpha_M^* = \alpha_P + \frac{\rho_2 \sigma_2^p \sigma_2^v \sigma_1^{p^2} - \rho_1 \sigma_1^p \sigma_1^v \sigma_2^{p^2}}{c \left(\sigma_1^{p^2} \sigma_2^{p^2} + \sigma_1^{p^2} + \sigma_2^{p^2} \right) + k \sigma_1^{p^2} \sigma_2^{p^2}}. \quad (12)$$

For interior solutions, the manager's preferred resource allocation α_M is increasing in ρ_2 and decreasing in ρ_1 . If $\rho_2 \sigma_2^p \sigma_2^v \sigma_1^{p^2} > \rho_1 \sigma_1^p \sigma_1^v \sigma_2^{p^2}$ (i.e. $\alpha_M > \alpha_P$), then α_M is decreasing in c and k .

For $\alpha_P \in (\underline{\alpha_M}, \overline{\alpha_M})$, the firm appoints a manager with $\alpha_M > \alpha_P$ if and only if:

$$\rho_2 \frac{\sigma_2^v}{\sigma_2^p} > \rho_1 \frac{\sigma_1^v}{\sigma_1^p}.$$

The type of manager hired depends on the firm's performance measurement system. The firm appoints a manager biased towards dimension 1 ($\alpha_M > \alpha_P$) if it is relatively poorly measured compared to dimension 2, i.e. ρ_2 is sufficiently high or ρ_1 is sufficiently low. This allows it to provide strong incentives on the well-measured dimension while relying on the manager's bias to mitigate excessive neglect of the poorly-measured dimension.

For example, with normalized volatilities $\sigma_1^v = \sigma_1^p = \sigma_2^v = \sigma_2^p = 1$, $\alpha_M > \alpha_P$ if and only if $\rho_2 > \rho_1$. Intuitively, incentives on dimension 2 are higher-powered to make the manager more responsive to his private information. To partly offset the effect of stronger incentives on resource allocation, the principal appoints a manager who is biased towards dimension 1. More generally, the optimal bias will be greater when output volatility on both dimensions is scaled

⁵In section 4, where we take into account the scarcity of any given managerial type, the firm does not choose α_M , so corner solutions do not arise.

up by the same amount, i.e., when the manager's private information matters more. The bias is decreasing in c and k . When the manager's deviation cost c is higher, then he is less responsive to incentives, reducing the need for bias to correct the distortion caused by incentives. When the firm's deviation cost k is higher, it wants a less biased manager because bias shifts him away from the firm's preferred allocation.

One application is to the case in which the first dimension is impact and the second is profitability. Assume that profitability is well-measured due to accounting standards, and that impact is poorly measured because it is difficult to quantify and the relevant measures of impact are inconsistent across firms. We again have $\rho_2 > \rho_1$ and so the firm optimally hires a manager biased toward impact, i.e. $\alpha_M > \alpha_P$. This further reduces the sensitivity of pay to impact (see equation (10)) and increases the sensitivity of pay to profitability (see equation (11)).

A second application is to product design, where the two dimensions are high quality and low cost. Cost is measured more accurately than quality, so ρ_2 is high and ρ_1 is low. The firm then optimally hires a product manager biased towards quality, and gives stronger incentives on the cost dimension.

A special case of Proposition 1 is where the second term is zero, and so the firm hires a congruent manager ($\alpha_M = \alpha_P$). Corollary 1 gives sufficient conditions for this to be the case.

Corollary 1 *The principal appoints a congruent manager in any of the following cases:*

- (i) *The manager has no private information on how the allocation affects output, i.e. $\sigma_1^v = \sigma_2^v = 0$;*
- (ii) *Performance measures are unrelated to output, i.e. $\rho_1 = \rho_2 = 0$;*
- (iii) *Performance measures are perfectly related to output, i.e. $\rho_1 = \rho_2 = 1$ and $\sigma_i^p = \sigma_i^v$ for $i = 1, 2$;*
- (iv) *The manager has no private information on how the allocation affects performance on one dimension, i.e. $\sigma_1^p = 0$ or $\sigma_2^p = 0$.*
- (v) *Both dimensions have the same performance measurement and volatility, i.e. $\rho_1 = \rho_2$, $\sigma_1^v = \sigma_2^v$ and $\sigma_1^p = \sigma_2^p$.*

In case (i), $\sigma_1^v = \sigma_2^v = 0$ and the manager has no private information on productivity. Then, the principal's goal is to implement her preferred allocation, $a = \alpha_P$, rather than to incentivize the manager to use his information. This can be achieved by hiring a congruent manager ($\alpha_M = \alpha_P$) with no incentives ($b_1 = b_2 = 0$).

In case (ii), $\rho_1 = \rho_2 = 0$ and performance measures are pure noise. Thus, optimal incentives are zero. The principal instead achieves her objective by hiring a congruent manager ($\alpha_M = \alpha_P$).

In case (iii), performance equals output and so the principal can achieve first best by incentivizing performance. However, the manager's preferences still matter because they affect

his disutility from deviation and thus required compensation. The firm thus hires a congruent manager.⁶

In case (iv), suppose $\sigma_1^p = 0$ (the argument is symmetric if $\sigma_2^p = 0$). The manager then has no private information on how the allocation affects measured performance on dimension 1, so incentives on that dimension do not induce the use of private information. Instead, they serve only to offset the distortion caused by incentives on dimension 2. The firm thus hires a congruent manager, $\alpha_M = \alpha_P$, and offers balanced incentives, $b_1 = b_2$ at a level driven by the measurement technology on dimension 2 (see Lemma 2 with $\sigma_1^p = 0$ and $\alpha_M = \alpha_P$).

In case (v), the two dimensions are symmetric. We then have $\alpha_M = \alpha_P$ and $b_1 = b_2$. An increase in α_P only affects the type of manager hired (α_M increases to the same extent) but not incentives. If performance is symmetrically measured, the principal wants to offer symmetric incentives to avoid any distortion, and instead implements her preferred allocation by appointing a congruent manager.

By choosing a manager who is biased towards the worse-measured dimension, the firm can provide stronger incentives for the better-measured dimension. Lemma 3 shows how the “incentive gap”, i.e. the difference between the incentives on each dimension, depends on the bias of the manager hired. We define $\Delta \equiv \alpha_M - \alpha_P$ as the manager’s bias for dimension 1, which we sometimes refer to simply as “bias.”

Lemma 3 *The incentive gap is given by:*

$$b_2(\Delta) - b_1(\Delta) = \frac{c \left[\rho_2 \sigma_2^p \sigma_2^v \sigma_1^{p2} - \rho_1 \sigma_1^p \sigma_1^v \sigma_2^{p2} + k(\sigma_1^{p2} + \sigma_2^{p2})\Delta \right]}{(c + k) \left(\sigma_1^{p2} \sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2} \right)}.$$

The incentive gap, $b_2(\Delta) - b_1(\Delta)$, is increasing in the manager’s bias, Δ . If the firm hired a congruent manager, then $\Delta = 0$. In this case, the incentive gap is positive if and only if $\rho_2 \sigma_2^p \sigma_2^v \sigma_1^{p2} > \rho_1 \sigma_1^p \sigma_1^v \sigma_2^{p2}$, in which case a manager biased towards dimension 1 is optimal (Proposition 1). As a result, when the firm hires the optimal manager, $\Delta > 0$ and the incentive gap widens. Intuitively, hiring a biased manager allows incentives to respond more strongly to differences in performance measurement.

Proposition 2 studies the benefit to the principal from hiring the optimal manager rather than a congruent manager, i.e. the value of a non-value-maximizing manager. Let Δ^* be the optimal bias, i.e. Δ when α_M is as in Proposition 1.

⁶Interestingly, this is no longer true if performance measures are perfectly informative ($\rho_1 = \rho_2 = 1$) but do not have the same volatility as output. Intuitively, in this case, pay-for-performance optimally varies across performance measures, so that it is generally suboptimal to hire a perfectly congruent agent.

Proposition 2 *The principal's gain from hiring the optimal manager rather than a congruent manager is given by:*

$$\Pi^* - \Pi_0 = \frac{k(\rho_2\sigma_2^p\sigma_2^v\sigma_1^{p2} - \rho_1\sigma_1^p\sigma_1^v\sigma_2^{p2})^2}{2(c+k)GD} \geq 0. \quad (13)$$

where $D \equiv \sigma_1^{p2}\sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2}$ and $G \equiv cD + k\sigma_1^{p2}\sigma_2^{p2}$. If $\rho_2\frac{\sigma_2^v}{\sigma_2^p} > \rho_1\frac{\sigma_1^v}{\sigma_1^p}$, this gain is increasing in ρ_2 and σ_2^v and decreasing in ρ_1 , σ_1^v , and c . The gain can also be expressed as:

$$\Pi^* - \Pi_0 = \frac{(c+k)GD}{2c^2k\sigma_1^{p4}}(b_2(\Delta^*) - b_2(0))^2. \quad (14)$$

The gain is zero if and only if $\rho_1\sigma_1^v/\sigma_1^p = \rho_2\sigma_2^v/\sigma_2^p$, as then $\alpha_M = \alpha_P$ from Proposition 1: congruence is optimal. If $\rho_2\sigma_2^p\sigma_2^v\sigma_1^{p2} > \rho_1\sigma_1^p\sigma_1^v\sigma_2^{p2}$, then the firm optimally hires a manager biased towards dimension 1, and the gain is increasing in $(b_2(\Delta^*) - b_2(0))$, i.e. the increase in incentives for dimension 2 as a result of hiring the optimal manager rather than a congruent manager.

The comparative statics are as follows. Without loss of generality, we consider the case in which dimension 2 is sufficiently better measured than dimension 1, i.e. $\rho_2\sigma_2^p\sigma_2^v\sigma_1^{p2} - \rho_1\sigma_1^p\sigma_1^v\sigma_2^{p2} > 0$. The gain is increasing in ρ_2 and σ_2^v . When dimension 2 is better measured, or the manager's private information on dimension 2 is more valuable, the firm wishes to give stronger incentives for dimension 2. This increases the value of the manager's bias towards dimension 1 as a counterweight. In contrast, the gain is decreasing in ρ_1 and σ_1^v : when the firm wishes to give stronger incentives for dimension 1, there is less value from a manager with intrinsic preferences for dimension 1. The gain is decreasing in the deviation cost c . A higher cost makes the manager less responsive to financial incentives. This reduces the distortion caused by asymmetric incentives and thus the need for correction through bias.

The value of a biased manager has additional applications beyond our one-period model. First, consider a firm with an incumbent congruent manager that faces a switching cost of replacing him with the optimal manager. The value of such a switch is higher when ρ_2 and σ_2^v are larger, and ρ_1 , σ_1^v , and c are smaller. Second, firms benefit from a broad managerial pipeline (high $\alpha_M - \overline{\alpha_M}$) rather than only candidates with $\alpha_M = \alpha_P$. Performance measurement technologies and the value of private information may change over time, making different manager types optimal in different circumstances. This is relevant for a company's succession planning: it is valuable to develop a diversity of potential successors rather than only those in the same mould as the CEO. It is also relevant for business schools who train the next generation of managers: there is value to developing executives with a range of preferences, including preferences for impact rather than only shareholder value.

In summary, the precision of performance measurement on each dimension determines which manager is hired and the sensitivity of pay to performance. When performance is a good measure of output on dimension i , the manager should be highly responsive to his signal, ϵ_i^p ,

and so incentives b_i should be high. To avoid excessive resource allocation to dimension i at the expense of dimension j , the principal optimally appoints a manager biased towards the latter.

4 Matching in competitive markets

We now extend the model to a continuum of firms that differ in their performance measurement, while retaining a continuum of managers with heterogeneous preferences. This analysis allows us to study the matching between firms and managers, as well as which managers receive the highest rents (in the single-firm model, the manager receives his reservation wage). We also show that, unlike the single-firm model, managers biased towards one dimension need not receive weaker incentives for that dimension, because they may be matched with firms that can measure it more precisely.

4.1 Setup

There is a continuum of firms that differ only in the quality of their measurement of dimension 1, ρ_1 : the type of firm i , ρ_1^i , is distributed on $[\underline{\rho}, \bar{\rho}]$ with CDF F_ρ . The quality of measurement of dimension 2, ρ_2 , is common across firms. This allows firms to be indexed by a single parameter, ρ_1^i . One application is where dimension 2 is profit, which is measured relatively uniformly due to accounting standards, and dimension 1 is impact, the measurement of which varies substantially depending on the nature of impact activities. Importantly, we do not make any assumptions on ρ_2 compared to $\underline{\rho}$ and $\bar{\rho}$, i.e. we do not assume that dimension 2 is better or worse measured than dimension 1. Where there is no confusion, we will use “measurement” or “measurement quality” to refer to the measurement of dimension 1, ρ_1^i , since ρ_2 is common to all firms, and “preference” for the manager’s preference for dimension 1.

Since volatilities are common across firms, we normalize them to 1 ($\sigma_1^v = \sigma_2^v = \sigma_1^p = \sigma_2^p = 1$) for simplicity. As before, managers differ only in their preferred resource allocations α_M : the type of manager i , α_M^i , is distributed on $[\underline{\alpha}_M, \bar{\alpha}_M]$ with CDF F_{α_M} .

We consider a matching model with equal masses of firms and managers and transferable utility. Let $Y(\rho_1, \rho_2; \alpha_M)$ denote the total surplus from a firm-manager match – the firm’s objective function plus the manager’s expected utility. It is given by:

$$Y(\rho_1, \rho_2; \alpha_M) = \mathbb{E} \left[V_1 + V_2 - \frac{k}{2}(a - \alpha_P)^2 - \frac{c}{2}(a - \alpha_M)^2 \right] \quad (15)$$

when the firm optimally chooses (b_1, b_2) and the manager optimally chooses a .

Since the model has transferable utility – the fixed wage S freely transfers surplus between the firm and the manager – the firm maximizes its own payoff by maximizing joint surplus Y

and then setting S to give the manager his competitive equilibrium wage. An equilibrium of the model is therefore such that: (i) all firm-manager matches are stable, i.e., no firm and manager who are not matched with each other can both improve their payoffs by matching; (ii) all firms offer the contract (b_1, b_2) that maximizes the joint surplus $Y(\rho_1, \rho_2; \alpha_M)$, with the fixed wage S set to give the manager his competitive equilibrium payoff; (iii) all managers choose the optimal resource allocation given their contracts.

4.2 Matching equilibrium

We start by showing that the matching assignment depends on the cross-derivative of the surplus with respect to the two matching variables (Becker (1973), Legros and Newman (2007)).

Lemma 4 *The surplus of a match $Y(\rho_1, \rho_2, \alpha_M)$ is submodular in ρ_1 and α_M :*

$$\frac{\partial^2 Y(\rho_1, \rho_2, \alpha_M)}{\partial \rho_1 \partial \alpha_M} < 0. \quad (16)$$

The negative cross-derivative means that a manager with stronger preferences for dimension 1 (higher α_M) generates more surplus at a firm with weaker measurement of dimension 1 (lower ρ_1). Intuitively, strong preferences and weak measurement are complements, because the former compensate for the latter. Conversely, a firm with good measurement of dimension 1 does not need a manager biased towards that dimension – it can provide strong incentives directly – and is better off hiring a manager focused on dimension 2 (low α_M) who can be redirected toward dimension 1 through financial incentives.

Submodularity in a matching model with transferable utility leads to negative assortative matching, as shown in Proposition 3:

Proposition 3 *The matching equilibrium involves negative assortative matching: firms with high ρ_1^i are matched with managers with low α_M^i . The matching function $\rho_1(\alpha_M)$, which gives the type ρ_1^i of the firm matched with a manager with type α_M^i , is strictly decreasing and given by:*

$$\rho_1(\alpha_M^i) = F_\rho^{-1}(1 - F_{\alpha_M}(\alpha_M^i)). \quad (17)$$

Firms with the best measurement of dimension 1 (highest ρ_1) hire the managers with weakest preferences for that dimension (lowest α_M), because of the complementarity shown in Lemma 4.

4.3 Equilibrium rents

We now characterize the rents earned by different managers in the matching equilibrium. The rent $U(\alpha_M)$ of a manager with type α_M is defined as:

$$U(\alpha_M) \equiv \mathbb{E}\left[W(a, \epsilon) - \frac{c}{2}(a - \alpha_M)^2 \middle| \alpha_M\right] - \bar{U} \quad (18)$$

where a is optimally chosen.

Proposition 4 *The marginal rent of a manager with type α_M is:*

$$U'(\alpha_M) = \frac{k}{3(c+k)} (\rho_2 - \rho_1(\alpha_M) - (3c+k)(\alpha_M - \alpha_P)), \quad (19)$$

and the rent schedule is:

$$U(\alpha_M) = U(\underline{\alpha}_M) + \int_{\underline{\alpha}_M}^{\alpha_M} \left\{ \frac{k(\rho_2 - \rho_1(x))}{3(c+k)} - \frac{k(3c+k)}{3(c+k)}(x - \alpha_P) \right\} dx. \quad (20)$$

The marginal rent is higher if the firm has worse measurement for dimension 1 (lower ρ_1), which increases the value of managers with stronger preferences for dimension 1. It is also higher if dimension 2 is better measured (higher ρ_2): a manager with stronger preferences for dimension 1 is more valuable when higher incentives can be provided on dimension 2. Ceteris paribus, the rent is lower if the manager is more biased, because of the higher deadweight loss.

The rent function $U(\alpha_M)$ is generally non-monotonic: agents with extreme biases (very high or very low α_M) may earn lower rents than agents with intermediate biases, depending on the distribution of firm types and model parameters. We can derive a closed-form characterization when firms' measurement ρ_1 and managers' preferences α_M are uniformly distributed.

Corollary 2 *Assume $\rho_1 \sim \mathcal{U}[\underline{\rho}, \bar{\rho}]$ and $\alpha_M \sim \mathcal{U}[\underline{\alpha}_M, \bar{\alpha}_M]$.*

(i) If $\alpha_P \in (\underline{\alpha}_M, \bar{\alpha}_M)$, $\rho_2 \in (\underline{\rho}, \bar{\rho})$, and $\frac{\bar{\rho} - \underline{\rho}}{\bar{\alpha}_M - \underline{\alpha}_M} < 3c + k$, then the rent function $U(\alpha_M)$ is single-peaked, with a unique interior maximizer $\hat{\alpha}_M$ satisfying

$$\hat{\alpha}_M = \alpha_P + \frac{\rho_2 - \rho_1(\hat{\alpha}_M)}{3c + k}, \quad (21)$$

where $\rho_1(\alpha_M) = \bar{\rho} - \frac{\alpha_M - \underline{\alpha}_M}{\bar{\alpha}_M - \underline{\alpha}_M}(\bar{\rho} - \underline{\rho})$.

(ii) $\hat{\alpha}_M > \alpha_P$ if and only if $\rho_2 > \rho_1(\hat{\alpha}_M)$; in particular, if $\rho_2 > \bar{\rho} - \frac{\bar{\rho} - \underline{\rho}}{\bar{\alpha}_M - \underline{\alpha}_M}(\alpha_P - \underline{\alpha}_M)$, then the interior maximizer (if it exists) is such that $\hat{\alpha}_M > \alpha_P$.

Intuitively, managers who earn the highest rents are those whose bias complements the average measurement technology in the economy. $\rho_1(\hat{\alpha}_M)$ is the technology of the firm that hires

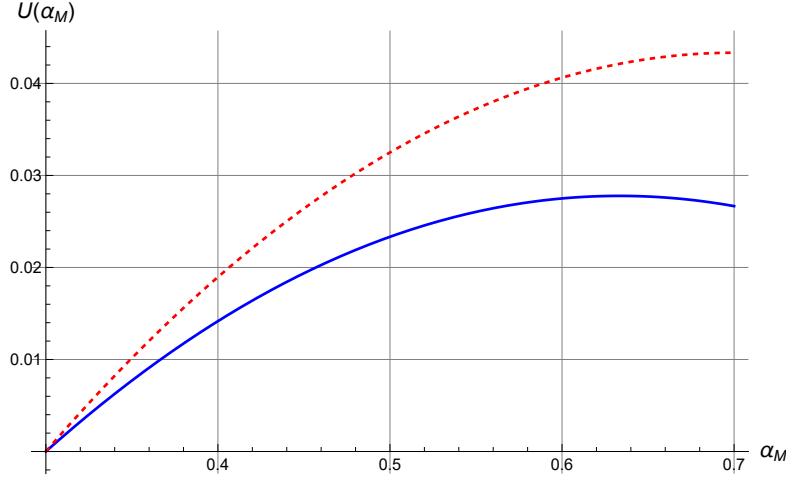


Figure 1: Comparison of rent schedules for managers with different preferences α_M under different measurement quality for dimension 1. The solid blue line is the agency rent in Example 1 with $\rho_1 \in [0.3, 0.7]$. The dashed red line is the agency rent in Example 2 with $\rho_1 \in [0.1, 0.4]$.

the rent-maximizing manager $\hat{\alpha}_M$. When $\rho_2 > \rho_1(\hat{\alpha}_M)$, i.e. when dimension 2 is sufficiently well-measured compared to dimension 1, the managers who earn the highest rents are more biased toward dimension 1 than firms. Their bias is valuable because it substitutes for poor measurement. However, if the bias is too large, the manager creates less value and earns fewer rents.

We now consider two examples that allow us to solve numerically for the manager's rent as a function of his preferences.

Example 1 $\rho_1 \sim \mathcal{U}[0.3, 0.7]$, $\rho_2 = 0.9$, $\alpha_M \sim \mathcal{U}[0.3, 0.7]$, $\alpha_P = 0.5$, $c = k = 1$. The highest rent is at $\hat{\alpha}_M = 0.63$, which is above $\alpha_P = 0.5$: the manager is biased towards dimension 1. Thus, managers with stronger preferences for dimension 1 earn the highest rents when it is measured much worse than dimension 2.

Example 2 $\rho_1 \sim \mathcal{U}[0.1, 0.4]$ while all other parameters are unchanged. Measurement of dimension 1 is worse than in Example 1. The highest rent is now at $\hat{\alpha}_M = 0.7$: it is earned by managers with even stronger preferences for dimension 1 than in Example 1.

Figure 1 depicts the two rent schedules. As measurement worsens from Example 1 to Example 2, the demand for managers with stronger preferences for dimension 1 increases. This raises rents overall, as the least valuable manager type must still be willing to participate, with a rent of zero. Moreover, when measurement is moderately poor (Example 1), the highest rent is earned by a manager with an intermediate bias, whereas when measurement is very poor (Example 2), the highest rent accrues to the manager with the most extreme bias.

4.4 Equilibrium incentives

Having solved for the firm a manager works for and the level of rents he receives, we finally study the incentives he is given. From Lemma 2, incentives are given by:

$$b_1 = \frac{c(2\rho_1 + \rho_2 - k\Delta)}{3(c+k)} \quad (22)$$

$$b_2 = \frac{c(\rho_1 + 2\rho_2 + k\Delta)}{3(c+k)}. \quad (23)$$

To study how manager type affects incentives, we take the total derivative with respect to α_M , which yields:

$$\begin{aligned} \frac{db_1}{d\alpha_M} &= \frac{c}{3(c+k)} [-k + 2\rho'_1(\alpha_M)] \\ \frac{db_2}{d\alpha_M} &= \frac{c}{3(c+k)} [k + \rho'_1(\alpha_M)]. \end{aligned}$$

The effect of manager type on incentives can thus be decomposed into two components. The first is the *direct effect* holding the matched firm constant, which is given by the first term in square brackets. As in the single-firm model, a manager biased towards dimension 1 receives weaker incentives on dimension 1 and stronger incentives on dimension 2. The second is the *sorting effect*: the manager will be matched with a firm that measures dimension 1 less accurately: $\rho'_1(\alpha_M) < 0$ due to negative assortative matching. Such a firm provides weaker incentives on dimension 1, and also on dimension 2 due to the spillover effect in Lemma 2. This reinforces the direct effect for b_1 (both terms in square brackets are negative) but counteracts the direct effect for b_2 . The overall effect is given in Proposition 5:

Proposition 5 *At any α_M , the equilibrium incentive $b_1(\alpha_M)$ is strictly decreasing in α_M . The equilibrium incentive $b_2(\alpha_M)$ is locally increasing in α_M if $k > -\rho'_1(\alpha_M)$, and locally decreasing if $k < -\rho'_1(\alpha_M)$.*

Consider two managers with types $\alpha_M^H > \alpha_M^L$. Manager H always receives weaker incentives for dimension 1 than manager L . He receives stronger incentives for dimension 2 if $k > -\rho'_1(\alpha_M)$ holds for all $\alpha_M \in [\alpha_M^L, \alpha_M^H]$, and weaker incentives if $k < -\rho'_1(\alpha_M)$ holds for all $\alpha_M \in [\alpha_M^L, \alpha_M^H]$.

The direct effect is increasing in the firm's adjustment cost k . A biased manager (high Δ) allocates too many resources to dimension 1. This is more costly to the firm when k is high, so it responds by providing strong incentives for dimension 2: in equation (23), b_2 is increasing in $k\Delta$. The sorting effect is increasing in the slope of the matching function $-\rho'_1(\alpha_M)$. From Proposition 3, we have

$$-\rho'_1(\alpha_M) = \frac{f_{\alpha_M}(\alpha_M)}{f_\rho(\rho_1(\alpha_M))}.$$

This is large when variation in firm measurement is high relative to variation in manager preferences. When ρ_1 and α_M are both uniformly distributed, this becomes:

$$-\rho'_1(\alpha_M) = \frac{\bar{\rho} - \underline{\rho}}{\bar{\alpha}_M - \underline{\alpha}_M}.$$

The slope is then constant, so the comparison between any two types reduces to $k \gtrless (\bar{\rho} - \underline{\rho})/(\bar{\alpha}_M - \underline{\alpha}_M)$. If $\bar{\rho} - \underline{\rho}$ is large and $\bar{\alpha}_M - \underline{\alpha}_M$ is small, managers with only slightly different preferences will be matched to firms with substantially different measurement, leading to a large sorting effect.

Proposition 5 shows how endogenous matching can either preserve or reverse the single-firm relation between managerial bias and incentives in Lemma 2. If the direct effect dominates, managers with stronger preferences for dimension 1 receive higher incentives for dimension 2, as in the single-firm model – even in a market equilibrium with heterogeneous firms and managers. If the sorting effect dominates, the cross-sectional relation reverses. Managers with stronger preferences for dimension 1 are matched with firms that measure dimension 1 poorly, which weakens incentives on dimension 1 and, through the spillover effect in Lemma 2, also weakens incentives on dimension 2. This prediction of weak incentives on both dimensions is generated by matching rather than by the standard multitasking channel. Conversely, managers with weaker preferences for dimension 1 are matched with firms that measure dimension 1 well and therefore receive stronger incentives on both dimensions.

Figure 2 illustrates these two regimes using the same parameter values as in Examples 1 and 2, varying only k . In Panel A, with $k = 2$, the direct effect dominates: managers with stronger preferences for dimension 1 receive higher incentives for dimension 2 and weaker incentives for dimension 1. In Panel B, with $k = 0.5$, the sorting effect dominates: they receive weaker incentives on both dimensions.

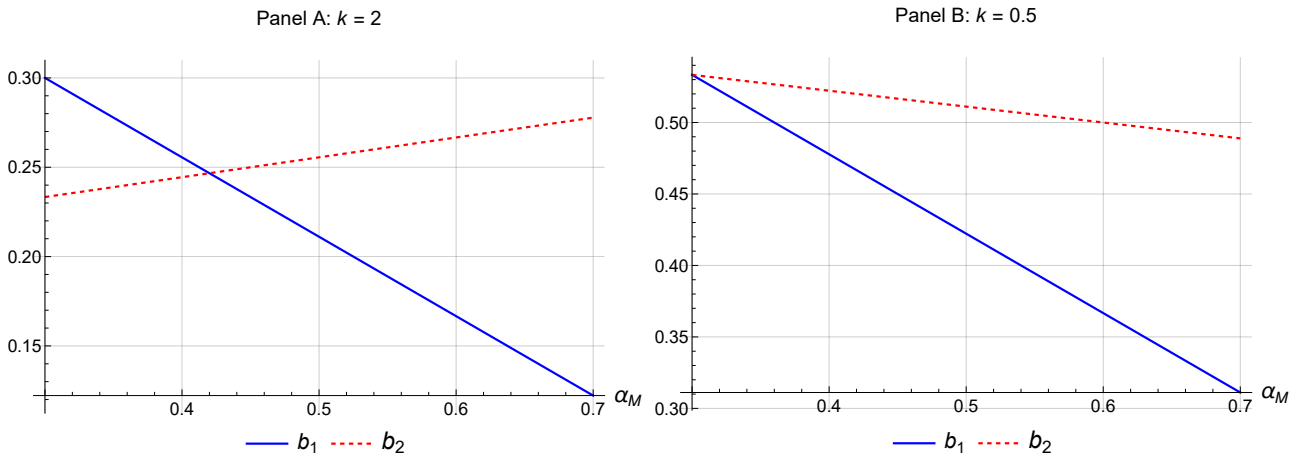


Figure 2: Equilibrium incentives b_1 and b_2 as a function of managerial preferences α_M .

5 Implications

The model yields implications for the joint distribution of measurement quality, managerial type, incentive contracts, and compensation. This section discusses these implications.

Implication 1 *Firms with worse measurement on a given dimension should appoint managers who are biased toward that dimension.*

Implication 1 concerns managerial selection. It holds in both the single-firm model of Section 3 and the matching model of Section 4.

Empirically, a stronger preference for impact might be reflected in donations, affiliations, and voting records (Christensen et al. (2017), Di Giuli and Kostovetsky (2014)). It can also be identified via exogenous shocks to impact preferences (e.g. Cronqvist and Yu (2017)). Better impact measurement would be reflected in a higher correlation between impact ratings or scores (Berg, Kölbel, and Rigobon (2022)), ES reporting standards, third-party verification, or anti-greenwashing enforcement.

Implication 2 *A manager who is more biased toward the poorly-measured dimension should receive weaker incentives on that dimension.*

Implication 2 concerns incentives. For example, more impact-motivated executives should have weaker impact incentives. This prediction holds in both the single-firm model and the matching equilibrium.

Implication 3 *A manager who is more biased toward the poorly-measured dimension should receive stronger incentives on the well-measured dimension if firms face large adjustment costs or if there is little cross-firm variation in measurement quality of the poorly-measured dimension. Otherwise, he receives weaker incentives on the well-measured dimension.*

Implication 3 emphasizes that the effects of performance measurement and managerial selection on incentives for the well-measured dimension act in opposite directions. The direct effect identified in the single-firm model of Section 3 is that a manager more biased towards impact will receive stronger profit incentives. The matching model of Section 4 identifies a potentially offsetting effect, which will dominate if there is enough variation in the quality of impact measurement across firms, or firms' adjustment costs are low.

Implication 4 *When firms or managers face large adjustment costs, the highest rent is obtained by a manager with an intermediate bias. Otherwise, the highest rent is obtained by a manager with an extreme bias.*

As already highlighted in Corollary 2, the rent can be either monotonic or nonmonotonic in the manager's type. When c and k are low, the value of compensating for poor measurement can dominate the deadweight loss from bias, so extreme biases may earn the highest rents. When c and k are high, this deadweight loss is larger and the rent-maximizing type is interior.

Implication 5 *Following a market-wide improvement in the quality of performance measurement on one dimension, firms will have a weaker preference for managers biased towards this dimension, and managers will receive stronger incentives for this dimension.*

If impact performance becomes better measured following mandatory disclosure regulation, firms will have a stronger preference for profit-oriented managers. They may still prefer managers with an impact bias (if impact remains worse-measured than profits) but their preferred bias will now be more moderate. In addition, a market-wide improvement in performance measurement on any dimension – which does not affect the ranking of firms on this dimension and therefore does not affect the equilibrium matching – results in higher-powered incentives on this dimension, so that managers are more responsive to their information.

6 Applications

We now consider several applications of our model, and discuss the empirical implications in each context.

6.1 Executive compensation

In some firms, shareholders have a purely financial objective: they seek to maximize long-term shareholder value. However, long-term shareholder value is non-contractible because the manager's consumption needs limit the extent to which he can be given restricted equity. Thus, he will be paid at least in part on the short- or medium-term stock price. Some business activities may improve long-term shareholder value but not be fully reflected in the stock price, such as customer satisfaction (Fornell et al., 2006), employee satisfaction (Edmans, 2011), and other intangible assets. This is in part because they are poorly measured, which also makes them difficult to incentivize directly. Our model suggests that this could be addressed by hiring managers with a preference for these activities, and incentivizing them for financial performance using equity incentives. Our single-firm model predicts that financial incentives will be stronger for managers with greater preferences for impact.⁷

⁷In addition, as in Baker (1992), these incentives will be stronger if the stock price is a better measure of firm value, for example due to improvements in stock market efficiency.

Another application is to venture capital. Early-stage firms often have measurable short-term metrics such as user growth or revenue, but their long-term value depends on poorly measured dimensions such as product quality, innovation, and organizational culture. Our model suggests that investors may optimally back founders who are intrinsically biased toward these harder-to-measure dimensions, while using strong financial incentives tied to observable metrics. This helps prevent excessive focus on short-term performance while preserving responsiveness to measurable outcomes.

While shareholders in the above examples care exclusively about long-term shareholder value, a third application concerns firms whose shareholders have non-financial objectives, such as societal impact (Hart and Zingales, 2017). Some proxy advisors assess a company’s executive compensation scheme in part on the relative strength of profit and ES incentives, and commentators often interpret strong ES (profit) incentives as an unambiguously positive (negative) commitment to a societal mission. Our model cautions against such interpretations. ES incentives may be an ineffective way to pursue impact given measurement difficulties; instead, hiring a manager with intrinsic preferences for impact may be more effective. Indeed, ES incentives should be optimally low if such a manager is hired. Similarly, an increase in ES incentives is only warranted following improvements in impact measurement, such as improved disclosure standards or third-party verification such as sustainability audits. Otherwise, such an increase would increase “hitting the target, missing the point” – improving measured rather than actual impact.

Finally, the model identifies the settings in which diversity in managerial preferences is valuable: when the relative quality of performance measurement varies across firms. In such environments, firms differ in the extent to which they can rely on incentives: firms with weak measurement of impact rely more on bias and thus prefer managers who place greater weight on impact, whereas firms with strong measurement rely more on incentives and prefer managers with weaker impact preferences. As a result, a heterogeneous pool of managerial types improves matching efficiency. This has implications for policies that homogenize or diversify the pool of available managers, such as MBA admissions and curricula, and for the value of cognitive diversity within organizations.

6.2 Universities

A university employs professors who engage in a range of activities. For simplicity, we will focus on just two (research and teaching), although the implications extend to other activities such as internal contribution and external visibility. These professors must decide how to allocate their time between teaching and research.

Some institutions place greater weight on the quantity of research output, which is relatively

easy to incentivize – for example, through bonuses or pay rises tied to the number of publications. If they primarily teach undergraduate students, who may be less able to assess teaching quality (for example, because these students prioritize test-taking skills over critical thinking or real-world relevance), common measures such as student evaluations are noisy and easily gamed. In such settings, explicit incentives for teaching may backfire. Instead, these institutions rely on selection and culture, hiring faculty with intrinsic motivation for teaching and fostering a “teaching culture” to sustain effort in this dimension.

Other institutions emphasize research quality and societal impact, which are difficult to measure, particularly in the short-term. If they primarily teach more experienced students, such as MBAs, Masters, and PhDs, who can better evaluate teaching, performance on the teaching dimension becomes more informative. These institutions can therefore rely more on explicit incentives for teaching, such as overload payments or executive education opportunities, while selecting faculty with strong intrinsic motivation for research.

In summary, the model predicts that universities will rely more on intrinsic motivation and culture in dimensions that are difficult to measure and prone to gaming, and use explicit incentives for better-measured dimensions.⁸

6.3 Healthcare

Medical doctors allocate their time across different activities. Some (“cure”), such as performing procedures, ordering tests, and prescribing drugs, generate measurable outputs and are often linked to compensation through fee-for-service payments or productivity targets. Others (“prevention”), such as giving advice and better understanding individual cases to avoid unnecessary procedures, affect longer-term outcomes that are difficult to measure. Since doctors are best placed to assess which activities are appropriate for a given patient, this creates a classic resource-allocation problem with private information.

Our model suggests that hospitals should hire doctors with an intrinsic preference for prevention, and complement this with incentives tied to cure. Concerns about overuse of revenue-generating procedures and underinvestment in prevention (e.g. Gawande (2009)) may reflect not the use of incentives per se, but a mismatch between incentives and the preferences of the doctors to whom they are applied. When doctors lack intrinsic motivation for patient-centered care, incentives tied to measurable outputs can lead to distortions. In such cases, the solution is not to reduce incentives but to hire doctors with strong preferences for prevention and create a patient-first culture. Only if this is not possible (e.g. because doctor preferences are relatively homogeneous or culture is difficult to change) should they reduce incentives for measured

⁸Alternatively, if all measures are easy to game, then incentives will be low-powered on all dimensions, as also predicted by other models.

performance.

7 Conclusion

This paper develops a new theory of managerial selection and incentive provision. We study a setting in which a manager with private information allocates resources across competing activities. When performance is imperfectly and asymmetrically measured, the principal optimally appoints a biased manager – one whose preferences differ from her own. In particular, the principal prefers a manager biased toward the more poorly measured dimension, as this bias mitigates distortions in resource allocation induced by incentive provision. As a result, changes in the measurement technology not only affect optimal contracts, but also the type of managerial preferences that are most valuable.

The model explains the coexistence of high-powered incentives on well-measured dimensions and low-powered incentives on poorly measured ones, even when these activities are substitutes. By selecting a biased manager, the principal can rely on the manager’s preferences to ensure resource allocation to poorly measured activities, while using strong incentives where they are most effective. This selection channel amplifies asymmetries in pay-for-performance across dimensions and highlights the importance of jointly determining compensation and hiring policies.

Our results build on the idea that an important part of education is to endow agents destined for careers in some field with preferences appropriate for this field (Akerlof and Kranton (2005), Van den Steen (2010)). Our contribution is to show that fostering a diversity of preferences can improve matching between workers and firms, and that the most valuable preferences to endow depend on the nature of performance measurement.

8 References

- Aghion, P. and Stein, J.C., 2008. Growth versus margins: Destabilizing consequences of giving the stock market what it wants. *Journal of Finance* 63, 1025-1058.
- Akerlof, G.A. and Kranton, R.E., 2005. Identity and the economics of organizations. *Journal of Economic Perspectives* 19, 9-32.
- Akerlof, G.A. and Kranton, R.E., 2000. Economics and identity. *Quarterly Journal of Economics* 115, 715-753.
- Baker, G.P., 1992. Incentive contracts and performance measurement. *Journal of Political Economy* 100, 598-614.
- Barron, D., Georgiadis, G. and Swinkels, J., 2020. Optimal contracts with a risk-taking agent. *Theoretical Economics* 15, 715-761.
- Bebchuk, L.A. and Tallarita, R., 2023. Will corporations deliver value to all stakeholders? *Vanderbilt Law Review* 75, 1031-1091.
- Becker, G. S. 1973. A theory of marriage: Part I. *Journal of Political Economy*, 81, 813-846.
- Ben-David, I. and Chinco, A., 2026. The max EPS paradigm for corporate finance. Working paper, National Bureau of Economic Research.
- Bénabou, R. and Tirole, J., 2016. Bonus culture: Competitive pay, screening, and multi-tasking. *Journal of Political Economy* 124, 305-370.
- Berg, F., Kölbel, J.F. and Rigobon, R., 2022. Aggregate confusion: The divergence of impact ratings. *Review of Finance* 26, 1315-1344.
- Besley, T. and Ghatak, M., 2005. Competition and incentives with motivated agents. *American Economic Review* 95, 616-636.
- Besley, T. and Ghatak, M., 2006. Sorting with motivated agents: implications for school competition and teacher incentives. *Journal of the European Economic Association* 4, 404-414.
- Besley, T. and Ghatak, M., 2018. Prosocial motivation and incentives. *Annual Review of Economics* 10, 411-438.
- Carlin, B.I. and Gervais, S., 2009. Work ethic, employment contracts, and firm value. *Journal of Finance* 64, 785-821.
- Cassar, L. and Meier, S., 2018. Nonmonetary incentives and the implications of work as a source of meaning. *Journal of Economic Perspectives* 32, 215-238.
- Christensen, D.M., Mikhail, M., Walther, B. and Wellman, L.A., 2017. From K Street to Wall Street: Politically connected analysts and stock recommendations. *The Accounting Review* 92, 1-28.
- Cronqvist, H. and Yu, F., 2017. Shaped by their daughters: Executives, female socialization, and corporate social responsibility. *Journal of Financial Economics* 126, 543-562.
- Di Giuli, A. and Kostovetsky, L., 2014. Are red or blue companies more likely to go green?

Politics and corporate social responsibility. *Journal of Financial Economics* 111, 158-180.

Edmans, A., 2011. Does the stock market fully value intangibles? Employee satisfaction and equity prices. *Journal of Financial Economics* 101, 621-640.

Edmans, A., Gosling, T. and Jenter, D., 2023. CEO compensation: Evidence from the field. *Journal of Financial Economics* 150, 103718.

Fehr, E. and Charness, G., 2025. Social preferences: fundamental characteristics and economic consequences. *Journal of Economic Literature* 63, 440-514.

Fischer, P. and Huddart, S., 2008. Optimal contracting with endogenous social norms. *American Economic Review* 98, 1459-1475.

Fornell, C., Mithas, S., Morgeson III, F.V. and Krishnan, M.S., 2006. Customer satisfaction and stock prices: High returns, low risk. *Journal of Marketing*, 70, 3-14.

Gawande, A., 2009. The cost conundrum. *Annals of Medicine: The Cost Conundrum: Reporting & Essays: The New Yorker*, June 1.

Hart, O. and Zingales, L., 2017. Companies should maximize shareholder welfare not market value. *Journal of Law, Finance, and Accounting*, 2, 247-275.

Holmström, B. and Milgrom, P., 1991. Multitask principal-agent analyses: Incentive contracts, asset ownership, and job design. *Journal of Law, Economics, and Organization* 7, 24-52.

Kerr, S., 1975. On the folly of rewarding A, while hoping for B. *Academy of Management Journal* 18, 769-783.

Legros P. and A. Newman 2007. Beauty is a beast, frog is a prince: Assortative matching with non-transferabilities. *Econometrica* 75, 1073-1102.

Prendergast, C., 2007. The motivation and bias of bureaucrats. *American Economic Review* 97, 180-196.

Prendergast, C., 2008. Intrinsic motivation and incentives. *American Economic Review* 98, 201-205.

Terviö, M., 2008. The Difference That CEOs Make: An Assignment Model Approach. *American Economic Review* 98, 642-668.

Van den Steen, E., 2010. On the origin of shared beliefs (and corporate culture). *RAND Journal of Economics* 41, 617-648.

Vladasel, T., Parker, S.C., Sloof, R. and Van Praag, M., 2024. Revenue drift, incentives, and effort allocation in social enterprises. *Journal of Economics & Management Strategy* 33, 630-651.

Wrzesniewski, A., McCauley, C., Rozin, P. and Schwartz, B., 1997. Jobs, careers, and callings: People's relations to their work. *Journal of Research in Personality* 31, 21-33.

9 Appendix

Proof of Lemma 1: The agent chooses a to maximize his objective function in equation (5). At the time of choosing his action, the agent observes ϵ_1^p and ϵ_2^p , so that his objective function can be rewritten as:

$$S + b_1 a \epsilon_1^p + b_2 (1 - a) \epsilon_2^p - \frac{c}{2} (a - \alpha_M)^2$$

This objective function is concave in a , so that the optimum is given by the first-order condition (FOC):

$$b_1 \epsilon_1^p - b_2 \epsilon_2^p - c (a - \alpha_M) = 0$$

Rearranging gives the expression in Lemma 1. ■

Proof of Lemma 2: Plugging equation (4) into equation (3), the objective function of the principal is:

$$\max_{b, S} \mathbb{E} \left[V_1(a, \epsilon) + V_2(a, \epsilon) - S - b_1 P_1(a, \epsilon) - b_2 P_2(a, \epsilon) - \frac{k}{2} (a - \alpha_P)^2 \right]$$

The participation constraint and incentive constraint are respectively:

$$\mathbb{E} \left[S + b_1 P_1(a, \epsilon) + b_2 P_2(a, \epsilon) - \frac{c}{2} (a - \alpha_M)^2 \right] \geq \bar{U} \quad (24)$$

$$b_1 \epsilon_1^p - b_2 \epsilon_2^p - c (a - \alpha_M) = 0 \quad \Leftrightarrow \quad a = \alpha_M + \frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \quad (25)$$

Substituting in equation (3) for the binding participation constraint in equation (24), the principal's optimization program rewrites as:

$$\max_b \mathbb{E} \left[V_1(a, \epsilon) + V_2(a, \epsilon) - \frac{c}{2} (a - \alpha_M)^2 - \bar{U} - \frac{k}{2} (a - \alpha_P)^2 \right]$$

This objective function can be rewritten as:

$$\begin{aligned} & \mathbb{E} \left[\left(\alpha_M + \frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \right) \epsilon_1^v + \left(1 - \alpha_M - \frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \right) \epsilon_2^v - \frac{c}{2} \left(\frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \right)^2 - \bar{U} \right. \\ & \left. - \frac{k}{2} \left(\alpha_M - \alpha_P + \frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \right)^2 \right] \end{aligned}$$

The FOC with respect to b_1 is:

$$\begin{aligned}
& \mathbb{E} \left[\frac{\epsilon_1^p}{c} \epsilon_1^v - \frac{\epsilon_1^p}{c} \epsilon_2^v - c \epsilon_1^p \left(\frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c^2} \right) - k \frac{\epsilon_1^p}{c} \left(\alpha_M - \alpha_P + \frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \right) \right] = 0 \\
& \Leftrightarrow \mathbb{E} [\epsilon_1^p \epsilon_1^v] - \mathbb{E} [\epsilon_1^p \epsilon_2^v] - b_1 \mathbb{E} [\epsilon_1^{p2}] + b_2 \mathbb{E} [\epsilon_1^p \epsilon_2^p] - k \left((\alpha_M - \alpha_P) \mathbb{E} [\epsilon_1^p] + \frac{b_1}{c} \mathbb{E} [\epsilon_1^{p2}] - \frac{b_2}{c} \mathbb{E} [\epsilon_1^p \epsilon_2^p] \right) = 0 \\
& \Leftrightarrow b_1 = \frac{\mathbb{E} [\epsilon_1^p \epsilon_1^v] - \mathbb{E} [\epsilon_1^p \epsilon_2^v] - k(\alpha_M - \alpha_P) \mathbb{E} [\epsilon_1^p] + b_2 \left(1 + \frac{k}{c}\right) \mathbb{E} [\epsilon_1^p \epsilon_2^p]}{\left(1 + \frac{k}{c}\right) \mathbb{E} [\epsilon_1^{p2}]} \\
& \Leftrightarrow b_1 = \frac{\rho_1 \sigma_1^v \sigma_1^p - k(\alpha_M - \alpha_P)}{\left(1 + \frac{k}{c}\right) (\sigma_1^{p2} + 1)} + b_2 \frac{1}{\sigma_1^{p2} + 1} \tag{26}
\end{aligned}$$

Likewise, the FOC with respect to b_2 is:

$$\begin{aligned}
& \mathbb{E} \left[-\frac{\epsilon_2^p}{c} \epsilon_1^v + \frac{\epsilon_2^p}{c} \epsilon_2^v + c \epsilon_2^p \left(\frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c^2} \right) + k \frac{\epsilon_2^p}{c} \left(\alpha_M - \alpha_P + \frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \right) \right] = 0 \\
& \Leftrightarrow b_2 = \frac{\rho_2 \sigma_2^v \sigma_2^p + k(\alpha_M - \alpha_P)}{\left(1 + \frac{k}{c}\right) (\sigma_2^{p2} + 1)} + b_1 \frac{1}{\sigma_2^{p2} + 1} \tag{27}
\end{aligned}$$

Solving equations (26) and (27) for b_1 and b_2 yields:

$$b_1 = \frac{c \left(\sigma_2^{p2} (-k(\alpha_M - \alpha_P)) + \rho_2 \sigma_2^p \sigma_2^v + \rho_1 \sigma_1^p \sigma_1^v (\sigma_2^{p2} + 1) \right)}{(c + k) (\sigma_1^{p2} \sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2})} \tag{28}$$

$$b_2 = \frac{c \left(\sigma_1^{p2} k(\alpha_M - \alpha_P) + \rho_1 \sigma_1^p \sigma_1^v + \rho_2 \sigma_2^p \sigma_2^v (\sigma_1^{p2} + 1) \right)}{(c + k) (\sigma_1^{p2} \sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2})} \tag{29}$$

■

Proof of Proposition 1: The principal's objective function is:

$$\begin{aligned}
& \mathbb{E} \left[\left(\alpha_M + \frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \right) \epsilon_1^v + \left(1 - \alpha_M - \frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \right) \epsilon_2^v \right. \\
& \quad \left. - \frac{c}{2} \left(\frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \right)^2 - \bar{U} - \frac{k}{2} \left(\alpha_M - \alpha_P + \frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \right)^2 \right] \\
& = \left(\alpha_M \mathbb{E} [\epsilon_1^v] + \frac{b_1 \mathbb{E} [\epsilon_1^p \epsilon_1^v] - b_2 \mathbb{E} [\epsilon_2^p \epsilon_1^v]}{c} \right) + \left((1 - \alpha_M) \mathbb{E} [\epsilon_2^v] - \frac{b_1 \mathbb{E} [\epsilon_1^p \epsilon_2^v] - b_2 \mathbb{E} [\epsilon_2^p \epsilon_2^v]}{c} \right) \\
& \quad - \frac{1}{2c} \mathbb{E} [(b_1 \epsilon_1^p - b_2 \epsilon_2^p)^2] - \bar{U} - \frac{k}{2} \mathbb{E} \left[\left(\alpha_M - \alpha_P + \frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \right)^2 \right] \\
& = 1 + \frac{1}{c} (b_1 (\mathbb{E} [\epsilon_1^p \epsilon_1^v] - 1) + b_2 (\mathbb{E} [\epsilon_2^p \epsilon_2^v] - 1)) \\
& \quad - \frac{1}{2c} (b_1^2 \mathbb{E} [\epsilon_1^{p2}] + b_2^2 \mathbb{E} [\epsilon_2^{p2}] - 2b_1 b_2) - \bar{U} - \frac{k}{2} \mathbb{E} \left[\left(\alpha_M - \alpha_P + \frac{b_1 \epsilon_1^p - b_2 \epsilon_2^p}{c} \right)^2 \right] \tag{30}
\end{aligned}$$

where b_1 and b_2 are described in Lemma 2. Since b_1 and b_2 are chosen optimally by the principal and are interior solutions, we use the envelope theorem that states that the total derivative of the objective function with respect to α_M is equal to the partial derivative holding b_1 and b_2 constant. This partial derivative is:

$$k \left(\alpha_P - \alpha_M + \frac{-b_1 + b_2}{c} \right) \tag{31}$$

Since the objective function of the principal is concave in α_M , setting this derivative to zero gives:

$$\alpha_M = \alpha_P + \frac{b_2 - b_1}{c} \tag{32}$$

Combining with the optimal values of b_1 and b_2 in Lemma 2 and reorganizing completes the proof.

The principal's objective function is strictly concave in α_M , so that it is maximized at a unique point on $[\underline{\alpha}_M, \overline{\alpha}_M]$, which is described by the FOC above if an interior solution exists in $[\underline{\alpha}_M, \overline{\alpha}_M]$, and by a corner solution otherwise. ■

Proof of Lemma 3:

From Lemma 2, the optimal pay-performance sensitivities for a given Δ are:

$$b_1(\Delta) = \frac{c \left[\rho_1 \sigma_1^p \sigma_1^v (\sigma_2^{p^2} + 1) + \rho_2 \sigma_2^p \sigma_2^v - k \sigma_2^{p^2} \Delta \right]}{(c + k) (\sigma_1^{p^2} \sigma_2^{p^2} + \sigma_1^{p^2} + \sigma_2^{p^2})}, \quad (33)$$

$$b_2(\Delta) = \frac{c \left[\rho_2 \sigma_2^p \sigma_2^v (\sigma_1^{p^2} + 1) + \rho_1 \sigma_1^p \sigma_1^v + k \sigma_1^{p^2} \Delta \right]}{(c + k) (\sigma_1^{p^2} \sigma_2^{p^2} + \sigma_1^{p^2} + \sigma_2^{p^2})}. \quad (34)$$

Since both expressions share the same denominator, taking the difference $b_2(\Delta) - b_1(\Delta)$ reduces to expanding the numerator:

$$\begin{aligned} & \rho_2 \sigma_2^p \sigma_2^v (\sigma_1^{p^2} + 1) + \rho_1 \sigma_1^p \sigma_1^v + k \sigma_1^{p^2} \Delta \\ & - \rho_1 \sigma_1^p \sigma_1^v (\sigma_2^{p^2} + 1) - \rho_2 \sigma_2^p \sigma_2^v + k \sigma_2^{p^2} \Delta \\ & = \rho_2 \sigma_2^p \sigma_2^v \sigma_1^{p^2} - \rho_1 \sigma_1^p \sigma_1^v \sigma_2^{p^2} + k (\sigma_1^{p^2} + \sigma_2^{p^2}) \Delta, \end{aligned} \quad (35)$$

where the level terms $\rho_1 \sigma_1^p \sigma_1^v$ and $\rho_2 \sigma_2^p \sigma_2^v$ cancel. This yields the incentive gap stated in Lemma 3:

$$b_2(\Delta) - b_1(\Delta) = \frac{c \left[\rho_2 \sigma_2^p \sigma_2^v \sigma_1^{p^2} - \rho_1 \sigma_1^p \sigma_1^v \sigma_2^{p^2} + k (\sigma_1^{p^2} + \sigma_2^{p^2}) \Delta \right]}{(c + k) (\sigma_1^{p^2} \sigma_2^{p^2} + \sigma_1^{p^2} + \sigma_2^{p^2})}.$$

The gap is increasing in Δ since the coefficient $ck(\sigma_1^{p^2} + \sigma_2^{p^2}) / [(c + k)(\sigma_1^{p^2} \sigma_2^{p^2} + \sigma_1^{p^2} + \sigma_2^{p^2})]$ is strictly positive. ■

Proof of Proposition 2: As derived in the proof of Proposition 1, the principal's payoff after substituting the manager's allocation from Lemma 1 and the binding participation constraint is:

$$\Pi(\alpha_M) = 1 - \bar{U} + \frac{b_1 Q_1 + b_2 Q_2}{c} - \frac{(c + k) \Sigma}{2c^2} - \frac{k \Delta^2}{2} - \frac{k \Delta (b_1 - b_2)}{c}, \quad (36)$$

where $Q_i \equiv \rho_i \sigma_i^p \sigma_i^v$, $\Sigma \equiv b_1^2 (\sigma_1^{p^2} + 1) + b_2^2 (\sigma_2^{p^2} + 1) - 2b_1 b_2$, and $b_1(\Delta)$, $b_2(\Delta)$ are as in Lemma 2.

Since $b_1(\Delta)$ and $b_2(\Delta)$ maximize (36) for each Δ , the FOCs $\frac{\partial \Pi}{\partial b_i} = 0$ hold. Using the envelope theorem and $\frac{\partial \Delta}{\partial \alpha_M} = 1$:

$$\frac{d\Pi}{d\alpha_M} = \frac{\partial \Pi}{\partial \Delta} \Big|_{b_1(\Delta), b_2(\Delta)} = -k\Delta - \frac{k(b_1(\Delta) - b_2(\Delta))}{c}. \quad (37)$$

Substituting the incentive gap from Lemma 3 into (37):

$$\begin{aligned}
\frac{d\Pi}{d\alpha_M} &= -k\Delta - \frac{k\left[\rho_1\sigma_1^p\sigma_1^v\sigma_2^{p2} - \rho_2\sigma_2^p\sigma_2^v\sigma_1^{p2} - k(\sigma_1^{p2} + \sigma_2^{p2})\Delta\right]}{(c+k)(\sigma_1^{p2}\sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2})} \\
&= -\frac{k}{(c+k)(\sigma_1^{p2}\sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2})} \\
&\quad \times \left[\left((c+k)(\sigma_1^{p2}\sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2}) - k(\sigma_1^{p2} + \sigma_2^{p2})\right)\Delta + \rho_1\sigma_1^p\sigma_1^v\sigma_2^{p2} - \rho_2\sigma_2^p\sigma_2^v\sigma_1^{p2}\right]. \quad (38)
\end{aligned}$$

The coefficient of Δ in the bracket simplifies:

$$\begin{aligned}
&(c+k)(\sigma_1^{p2}\sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2}) - k(\sigma_1^{p2} + \sigma_2^{p2}) \\
&= (c+k)\sigma_1^{p2}\sigma_2^{p2} + c(\sigma_1^{p2} + \sigma_2^{p2}) = G,
\end{aligned}$$

where the last equality uses the definition $G \equiv (c+k)\sigma_1^{p2}\sigma_2^{p2} + c(\sigma_1^{p2} + \sigma_2^{p2})$. Letting $\mathcal{A} \equiv \rho_2\sigma_2^p\sigma_2^v\sigma_1^{p2} - \rho_1\sigma_1^p\sigma_1^v\sigma_2^{p2}$, equation (38) becomes:

$$\frac{d\Pi}{d\alpha_M} = -\frac{k(G\Delta - \mathcal{A})}{(c+k)(\sigma_1^{p2}\sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2})}. \quad (39)$$

This is affine and strictly decreasing in Δ , confirming that Π is strictly concave in α_M . Setting (39) to zero gives $\Delta^* = \mathcal{A}/G$, recovering Proposition 1.

Integrating (39) from $\alpha_M^0 = \alpha_P$ (corresponding to $\Delta = 0$) to $\hat{\alpha}_M$ ($\Delta = \Delta^* = \mathcal{A}/G$):

$$\begin{aligned}
\Pi^* - \Pi_0 &= \int_0^{\Delta^*} \frac{-k(G\tilde{\Delta} - \mathcal{A})}{(c+k)(\sigma_1^{p2}\sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2})} d\tilde{\Delta} \\
&= \frac{k}{(c+k)(\sigma_1^{p2}\sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2})} \left[\mathcal{A}\Delta^* - \frac{G(\Delta^*)^2}{2} \right]. \quad (40)
\end{aligned}$$

Substituting $\Delta^* = \mathcal{A}/G$:

$$\mathcal{A}\Delta^* - \frac{G(\Delta^*)^2}{2} = \frac{\mathcal{A}^2}{G} - \frac{\mathcal{A}^2}{2G} = \frac{\mathcal{A}^2}{2G}.$$

Substituting into (40) and expanding $\mathcal{A}$:

$$\Pi^* - \Pi_0 = \frac{k\left(\rho_2\sigma_2^p\sigma_2^v\sigma_1^{p2} - \rho_1\sigma_1^p\sigma_1^v\sigma_2^{p2}\right)^2}{2(c+k)(\sigma_1^{p2}\sigma_2^{p2} + \sigma_1^{p2} + \sigma_2^{p2})G} \geq 0,$$

which is equation (13). The gain is zero if and only if $\mathcal{A} = 0$, i.e. $\rho_1\sigma_1^v/\sigma_1^p = \rho_2\sigma_2^v/\sigma_2^p$, the condition for $\hat{\alpha}_M = \alpha_P$ in Proposition 1.

From Lemma 2, the incentive rate on dimension 2 at the aligned benchmark ($\Delta = 0$) is:

$$b_2(0) = \frac{c \left[\rho_2 \sigma_2^p \sigma_2^v (\sigma_1^{p^2} + 1) + \rho_1 \sigma_1^p \sigma_1^v \right]}{(c+k)(\sigma_1^{p^2} \sigma_2^{p^2} + \sigma_1^{p^2} + \sigma_2^{p^2})}.$$

At the optimum ($\Delta = \Delta^* = \mathcal{A}/G$):

$$b_2(\Delta^*) = b_2(0) + \frac{ck\sigma_1^{p^2}\Delta^*}{(c+k)(\sigma_1^{p^2}\sigma_2^{p^2} + \sigma_1^{p^2} + \sigma_2^{p^2})} = b_2(0) + \frac{ck\sigma_1^{p^2}\mathcal{A}}{(c+k)(\sigma_1^{p^2}\sigma_2^{p^2} + \sigma_1^{p^2} + \sigma_2^{p^2})G}.$$

Therefore:

$$b_2(\Delta^*) - b_2(0) = \frac{ck\sigma_1^{p^2}\mathcal{A}}{(c+k)(\sigma_1^{p^2}\sigma_2^{p^2} + \sigma_1^{p^2} + \sigma_2^{p^2})G}. \quad (41)$$

Solving (41) for $\mathcal{A}$ and substituting into equation (13):

$$\begin{aligned} \Pi^* - \Pi_0 &= \frac{k}{2(c+k)(\sigma_1^{p^2}\sigma_2^{p^2} + \sigma_1^{p^2} + \sigma_2^{p^2})G} \cdot \frac{(c+k)^2(\sigma_1^{p^2}\sigma_2^{p^2} + \sigma_1^{p^2} + \sigma_2^{p^2})^2 G^2}{c^2 k^2 \sigma_1^{p^4}} (b_2(\Delta^*) - b_2(0))^2 \\ &= \frac{(c+k)(\sigma_1^{p^2}\sigma_2^{p^2} + \sigma_1^{p^2} + \sigma_2^{p^2})G}{2c^2 k \sigma_1^{p^4}} (b_2(\Delta^*) - b_2(0))^2, \end{aligned}$$

which is equation (14). ■

Proof of Corollary 1: The proofs of parts (i)-(iii) are special cases of the proof of Proposition 1.

The proof of part (iv) builds on the proof of Proposition 1. Evaluating the first derivative of the principal's objective function with respect to α_M in equation (31) at $\sigma_1^p = 0$ (substituting for the optimal values of b_1 and b_2) gives:

$$\frac{ck}{c+k} (\alpha_P - \alpha_M) \quad (42)$$

This expression has the same sign as $\alpha_P - \alpha_M$. Thus, if $\alpha_M < \alpha_P$, the principal would increase its objective function by increasing α_M ; if $\alpha_M > \alpha_P$, the principal would increase its objective function by decreasing α_M . Therefore, the principal's objective function is maximized at $\alpha_M = \alpha_P$. Plugging back into the optimal values of b_1 and b_2 yields $b_1 = b_2$. ■

Proof of Lemma 4: The surplus when the contract is chosen optimally is:

$$Y(\rho_1, \rho_2; \alpha_M) = 1 + \frac{b_1 \rho_1 + b_2 \rho_2}{c} - \frac{k\Delta^2}{2} - \Delta \frac{k(b_1 - b_2)}{c} - (b_1^2 + b_2^2 - b_1 b_2) \frac{c+k}{c^2} \quad (43)$$

where $\Delta \equiv \alpha_M - \alpha_P$, and with optimal incentive slopes:

$$b_1(\rho_1, \rho_2, \alpha_M) = \frac{c(-k\Delta + 2\rho_1 + \rho_2)}{3(c+k)} \quad (44)$$

$$b_2(\rho_1, \rho_2, \alpha_M) = \frac{c(k\Delta + \rho_1 + 2\rho_2)}{3(c+k)} \quad (45)$$

Substituting the optimal b_1, b_2 into equation (43) and simplifying yields the surplus in the form $Y = 1 + \frac{\frac{1}{3}(\rho_1^2 + \rho_1\rho_2 + \rho_2^2) + \frac{k}{3}(\rho_2 - \rho_1)\Delta - \frac{k(3c+k)}{6}\Delta^2}{c+k}$. Taking the derivative with respect to α_M (using $\frac{\partial \Delta}{\partial \alpha_M} = 1$):

$$\frac{\partial Y}{\partial \alpha_M} = \frac{k(\rho_2 - \rho_1)}{3(c+k)} - \frac{k(3c+k)}{3(c+k)}(\alpha_M - \alpha_P). \quad (46)$$

Now taking the derivative with respect to ρ_1 :

$$\frac{\partial^2 Y}{\partial \rho_1 \partial \alpha_M} = -\frac{k}{3(c+k)} < 0 \quad (47)$$

This establishes submodularity. ■

Proof of Proposition 4: The optimal incentive slopes are in equations (44) and (45), with $\Delta \equiv \alpha_M - \alpha_P$, and the joint surplus of a match is in equation (43). Substituting the optimal b_1, b_2 and simplifying yields:

$$Y(\rho_1, \rho_2; \alpha_M) = 1 + \frac{\frac{1}{3}(\rho_1^2 + \rho_1\rho_2 + \rho_2^2) + \frac{k}{3}(\rho_2 - \rho_1)\Delta - \frac{k(3c+k)}{6}\Delta^2}{c+k}.$$

Differentiating with respect to α_M (with $\frac{\partial \Delta}{\partial \alpha_M} = 1$):

$$\frac{\partial Y}{\partial \alpha_M} = \frac{k(\rho_2 - \rho_1)}{3(c+k)} - \frac{k(3c+k)}{3(c+k)}(\alpha_M - \alpha_P). \quad (48)$$

In a transferable-utility matching equilibrium, the marginal rent equals the marginal surplus generated by the agent (Terviö (2008)): when agent type α_M is matched with firm type $\rho_1(\alpha_M)$, we have $U'(\alpha_M) = \frac{\partial Y(\rho_1(\alpha_M), \rho_2; \alpha_M)}{\partial \alpha_M}$ (the partial with respect to agent type, evaluated at the matched firm). So:

$$U'(\alpha_M) = \frac{k}{3(c+k)} [\rho_2 - \rho_1(\alpha_M) - (3c+k)(\alpha_M - \alpha_P)]. \quad (49)$$

The rent schedule follows from the fundamental theorem of calculus: $U(\alpha_M) = U(\underline{\alpha_M}) + \int_{\underline{\alpha_M}}^{\alpha_M} U'(x)dx$. ■

Proof of Corollary 2: $U'(\alpha_M)$ is given in equation (49). Under uniform distributions, $F_{\alpha_M}(\alpha_M) = \frac{\alpha_M - \underline{\alpha}_M}{\bar{\alpha}_M - \underline{\alpha}_M}$ and $F_{\rho}^{-1}(q) = \underline{\rho} + q(\bar{\rho} - \underline{\rho})$, so that:

$$\rho_1(\alpha_M) = \bar{\rho} - \frac{\alpha_M - \underline{\alpha}_M}{\bar{\alpha}_M - \underline{\alpha}_M}(\bar{\rho} - \underline{\rho}), \quad \rho'_1(\alpha_M) = -\frac{\bar{\rho} - \underline{\rho}}{\bar{\alpha}_M - \underline{\alpha}_M} \equiv -m < 0.$$

Therefore:

$$U''(\alpha_M) = \frac{k}{3(c+k)} [-\rho'_1(\alpha_M) - (3c+k)] = \frac{k}{3(c+k)} (m - (3c+k)).$$

If $m < 3c+k$, then $U''(\alpha_M) < 0$ for all α_M , so $U(\cdot)$ is strictly concave and single-peaked with a unique interior maximizer. The FOC $U'(\hat{\alpha}_M) = 0$ gives

$$\rho_2 - \rho_1(\hat{\alpha}_M) - (3c+k)(\hat{\alpha}_M - \alpha_P) = 0 \quad \Leftrightarrow \quad \hat{\alpha}_M = \alpha_P + \frac{\rho_2 - \rho_1(\hat{\alpha}_M)}{3c+k}.$$

Substituting $\rho_1(\hat{\alpha}_M) = \bar{\rho} - m(\hat{\alpha}_M - \underline{\alpha}_M)$ yields :

$$\hat{\alpha}_M = \alpha_P + \frac{\rho_2 - \bar{\rho} + m(\alpha_P - \underline{\alpha}_M)}{(3c+k) - m}.$$

So part (i) holds.

For part (ii), $\hat{\alpha}_M > \alpha_P$ if and only if $\frac{\rho_2 - \rho_1(\hat{\alpha}_M)}{3c+k} > 0$, i.e., $\rho_2 > \rho_1(\hat{\alpha}_M)$. ■

Online Appendix

The Value of Non-Value-Maximizing Managers

Microfoundation for linear contracts

This Appendix adapts an argument from Barron, Georgiadis, and Swinkels (2020) to our setting.

We extend the model to account for the ex-post riskiness of the performance measure and the manager's ability to affect it. Suppose that the performance measure is as defined in the main text but is also noisy ex-post, so that $P_1(a, \epsilon) = a\epsilon_1^p + \varepsilon_1$ and $P_2(a, \epsilon) = (1 - a)\epsilon_2^p + \varepsilon_2$ where ε_1 and ε_2 are random variables with a mean of zero and a variance of one independent from other random variables that are unobservable.⁹ Suppose also that, conditional on a , the agent can implement arbitrary mean-preserving spreads and mean-preserving contractions of each performance measure distribution. Formally:

Assumption 1 *After choosing a resource allocation a , the agent can costlessly replace each performance measure P_i by any random variable $\tilde{P}_i$ such that:*

$$\mathbb{E}[\tilde{P}_i|a] = \mathbb{E}[P_i|a], \quad i = 1, 2.$$

A contract is defined as robust if every feasible joint distribution $(\tilde{P}_1, \tilde{P}_2)$ satisfying the mean constraints in Assumption 1 yields the same expected payment. Now consider a general contract $W(P_1, P_2)$ which can be a nonlinear function of performance measures.

Claim 1 *A contract $W(P_1, P_2) = S + \phi(P_1, P_2)$ is robust to the manipulation described in Assumption 1 if and only if:*

$$\phi(P_1, P_2) = b_0 + b_1 P_1 + b_2 P_2,$$

for some constants b_0, b_1, b_2 .

Proof of Claim 1: Given a resource allocation a , define $\mu_i \equiv \mathbb{E}[P_i|a]$ for $i = 1, 2$. Consider an arbitrary measurable contract $\phi : \mathbb{R}^2 \rightarrow \mathbb{R}$. The agent's problem at the manipulation stage is:

$$\sup_{(\tilde{P}_1, \tilde{P}_2)} \mathbb{E}[\phi(\tilde{P}_1, \tilde{P}_2)|a] \quad \text{s.t.} \quad \mathbb{E}[\tilde{P}_i|a] = \mu_i \quad \text{for } i = 1, 2.$$

⁹With universal risk neutrality, the presence of this pure risk does not change any of the relevant outcomes of the model.

Fix an arbitrary value $p_2 \in \mathbb{R}$ and define the univariate function of p_1 :

$$g_{p_2}(p_1) \equiv \phi(p_1, p_2).$$

Suppose that the function g_{p_2} is not affine in p_1 . Then there exist x, y , and $\lambda \in (0, 1)$ such that:

$$g_{p_2}(\lambda x + (1 - \lambda)y) \neq \lambda g_{p_2}(x) + (1 - \lambda)g_{p_2}(y).$$

There are two cases. First, if:

$$g_{p_2}(\lambda x + (1 - \lambda)y) < \lambda g_{p_2}(x) + (1 - \lambda)g_{p_2}(y), \quad (50)$$

then let $\tilde{P}_1$ take values x and y with probabilities λ and $1 - \lambda$, respectively, and set $P_1 = \lambda x + (1 - \lambda)y$ (i.e. P_1 is a constant). Then $\mathbb{E}[\tilde{P}_1|a] = \mathbb{E}[P_1|a]$, and

$$g_{p_2}(\mathbb{E}[\tilde{P}_1|a]) < \mathbb{E}[g_{p_2}(\tilde{P}_1)|a],$$

so the agent strictly increases his expected compensation with a mean-preserving spread. Second, if the inequality in equation (50) is reversed, a mean-preserving contraction yields the same conclusion.

Thus robustness requires that g_{p_2} be affine for every p_2 , i.e.

$$\phi(p_1, p_2) = \xi(p_2) + \beta(p_2)p_1$$

for some functions ξ and β . By symmetry, fixing p_1 and applying the same argument to $\phi(p_1, \cdot)$, robustness requires:

$$\phi(p_1, p_2) = \gamma(p_1) + \delta(p_1)p_2.$$

for some functions γ and δ . In sum, robustness requires:

$$\phi(p_1, p_2) = A + Bp_1 + Cp_2 + Dp_1p_2$$

for some constants A, B, C, D .

We now show by contradiction that a robust contract is such that $D = 0$. Suppose that $D \neq 0$. Fix the resource allocation a , and recall that $\mu_1 = \mathbb{E}[P_1|a]$ and $\mu_2 = \mathbb{E}[P_2|a]$. Pick any $\delta_1, \delta_2 > 0$ and consider two manipulations: (a) $(\tilde{P}_1, \tilde{P}_2)$ takes the values $(\mu_1 + \delta_1, \mu_2 + \delta_2)$ and $(\mu_1 - \delta_1, \mu_2 - \delta_2)$, each with probability $\frac{1}{2}$; (b) $(\tilde{P}_1, \tilde{P}_2)$ takes the values $(\mu_1 + \delta_1, \mu_2 - \delta_2)$ and $(\mu_1 - \delta_1, \mu_2 + \delta_2)$, each with probability $\frac{1}{2}$. Both have the same marginal distributions: $\tilde{P}_i \in \{\mu_i - \delta_i, \mu_i + \delta_i\}$ each with probability $\frac{1}{2}$, so $\mathbb{E}[\tilde{P}_i|a] = \mu_i$ for $i = 1, 2$ and both satisfy the

constraints in Assumption 1. However:

$$\mathbb{E}[\tilde{P}_1\tilde{P}_2|a] = \mu_1\mu_2 + \delta_1\delta_2 \quad \text{under (a),} \quad \mathbb{E}[\tilde{P}_1\tilde{P}_2|a] = \mu_1\mu_2 - \delta_1\delta_2 \quad \text{under (b).}$$

Since $\mathbb{E}[\phi(\tilde{P}_1, \tilde{P}_2)|a] = A + B\mu_1 + C\mu_2 + D\mathbb{E}[\tilde{P}_1\tilde{P}_2|a]$, the expected payment differs across (a) and (b) by $2D\delta_1\delta_2$, which is nonzero with $D \neq 0$, and these manipulations do not change the marginal means. This contradicts robustness, so $D = 0$.

If ϕ is linear, then for any admissible $(\tilde{P}_1, \tilde{P}_2)$, we can write:

$$\mathbb{E}[\phi(\tilde{P}_1, \tilde{P}_2)|a] = b_0 + b_1\mathbb{E}[\tilde{P}_1|a] + b_2\mathbb{E}[\tilde{P}_2|a],$$

which depends only on the means of the performance measures. Thus, no distributional manipulation can change the agent's expected pay. ■

about ECGI

The European Corporate Governance Institute has been established to improve *corporate governance through fostering independent scientific research and related activities*.

The ECGI will produce and disseminate high quality research while remaining close to the concerns and interests of corporate, financial and public policy makers. It will draw on the expertise of scholars from numerous countries and bring together a critical mass of expertise and interest to bear on this important subject.

The views expressed in this working paper are those of the authors, not those of the ECGI or its members.

ECGI Working Paper Series in Finance

Editorial Board

Editor	Nadya Malenko, Professor of Finance, Boston College
Consulting Editors	Renée Adams, Professor of Finance, University of Oxford Franklin Allen, Nippon Life Professor of Finance, Professor of Economics, The Wharton School of the University of Pennsylvania Julian Franks, Professor of Finance, London Business School Mireia Giné, Associate Professor, IESE Business School Marco Pagano, Professor of Economics, Facoltà di Economia Università di Napoli Federico II
Editorial Assistant	Asif Malik, Working Paper Series Manager

Electronic Access to the Working Paper Series

The full set of ECGI working papers can be accessed through the Institute's Web-site (www.ecgi.global/content/working-papers) or SSRN:

Finance Paper Series	http://www.ssrn.com/link/ECGI-Fin.html
Law Paper Series	http://www.ssrn.com/link/ECGI-Law.html