

Incentives to help or sabotage co-workers*

Evangelia Chalioti

Yale University[†]

Abstract

This paper studies compensation contracts and career concerns of co-workers when individual production depends on colleagues' effort and ability. We examine the conditions under which a manager who commits herself to a life-time salary path induces a worker to help or sabotage her colleagues. She may allow for little sabotage in order to decrease the provided insurance and motivate the workers to focus more on their own projects. If commitment is not feasible, workers now have incentives to help or sabotage due to career concerns. Such incentives arise for permanent and temporary workers, even though the latter will be paired to work with a different worker next period.

Keywords: career concerns, compensation contracts, sabotage incentives, relative performance

Jel codes: D83, J24, M54

*I am indebted to Lambros Pechlivanos for the very useful discussions on this topic. I am also grateful to Dan Bernhardt, Marco Celentani, Johannes Hörner, Patrick Rey for their comments and suggestions. I give special thanks to David Angenendt, Parimal Bag, V. Bhaskar, Roberto Burguet, Pierre Chaigneau, Gonzalo Cisternas, George Deltas, Laura Doval, Mehmet Ekmekci, George Georgiadis, Eiichiro Kazumori, Jorge Lemus, Vaiva Petrikaite, Mattias Polborn, Larry Samuelson, Konstantinos Serfes, Satoru Takahashi, Veikko Thiele, Caroline Thomas, Juuso Välimäki, Spyros Vassilakis and the participants at various conferences and meetings.

[†]Contact: Yale University, Department of Economics, 30 Hillhouse Ave., Room 35, New Haven, CT, 06511, USA. Email: evangelia.chalioti@yale.edu, Tel.: +1-217-607-3535

1 Introduction

In recent years, the re-engineering of R&D process in large corporations has shifted the organization of work towards various forms of teamworking and peer collaboration. The cross-functional and multi-disciplinary nature of research departments often entails that co-workers have different abilities that bring in the workplace. The degree to which each researcher contributes to a project and the number of years a researcher stays in a firm also vary. The U.S. Bureau of Labor Statistics shows that for 2016, the median tenure of engineers was 5.5 years, and that of mathematical and computer scientists was 4.4 years. Innovative firms also differ in their employment policies: they sign either long-term or short-term contracts with their workers, although a worker's research outcome also depends on the abilities of her colleagues with whom she interacts.¹ This paper studies how compensation contracts are shaped by managerial commitment and duration of employment when a worker's individual performance is influenced by her colleagues' effort and *ability*.²

We employ Holmström's (1982, 1999) career concerns framework where neither the workers nor the market know workers' innate abilities. We show that the manager who fully commits to a life-time income path may induce a low risk-averse worker to help her colleague, while she may motivate a high risk-averse worker even to sabotage. We argue that the manager may allow little sabotage in order to decrease the provided insurance and motivate the workers to focus on their own projects and increase their own outputs. If commitment is not feasible and short-term contracts are provided which are renegotiated in each period, we show that a worker now has incentives to sabotage her colleague because she wants to bias the learning process about her ability in her favor and built up her reputation. The manager now will use explicit compensation to decrease the worker's eagerness to sabotage. Reputational incentives to affect a colleague's performance arise even for temporary workers who will be paired with another worker in the next period. This happens in our model because a worker can manipulate market perceptions about her own ability also through her *current* colleagues' production.

In research labs, political parties, and modern corporations, the ability of a worker to be a good team player is considered an asset, but under certain conditions, allowing for 'little' sabotage may be optimal as a way of introducing competition among the workers. The sabotaged worker will also focus more on her own project, and improve her performance nonetheless. An article in The Telegraph (May 15, 2017, Sabotage? Complain? Have a duel?) discusses complains against Nick Robinson, co-presenter on

¹Empirical literature provides evidence that knowledge is transmitted among colleagues and skills diversity affects labor supply and productivity (Lazear (1999)). The benefits of peer interactions depend on whether the workers have distinct or identical knowledge and skills, indicating the degree of heterogeneity among the co-workers.

²The use of career concerns as an incentive device in the intra-firm relationships which may substitute explicit incentives due to compensation contracts is first discussed by Fama (1980) and elaborated by Holmström (1982). Gibbons & Murphy (1992) consider linear contracts and formalize this argument.

BBC Radio 4's Today program. Two of Robinson's four presenting colleagues – John Humphrys, Sarah Montague, Mishal Hussain and Justin Webb – complained that Robinson is beginning to dominate the program, using his decade-long tenure as BBC political editor as reason to demand he is in the chair for all high-profile news days. The article's suggestion of Robinson colleagues is to become more productive so as to regain control. As active sabotage, we consider, for example, the transmission of false information that can confuse or distract a colleague from the goals of her project.^{3,4} The benefit of designing contracts that induce sabotage is to decrease the variance of a worker's wage, and thus decrease the insurance that needs to be provided. The worker who experiences sabotage will also be motivated to focus more and work harder on her own project so as to undo the negative influence from her colleague and even increase her own production. Thus, from the managers' perspective, sabotage can be beneficial.

We consider a risk-neutral manager who appoints two risk-averse workers whose individual project outputs are observable and contractible, allowing the manager to treat workers separately through individual-based schemes (Itoh (1991), Itoh (1992), Auriol, Friebe & Pechlivanos (2002)).⁵ The incentive packages are derived in a linear principal-agent model (Holmström & Milgrom (1987), Gibbons & Murphy (1992)) and are based on explicit comparisons of co-workers' outputs. A worker's production function is linear in her "work" effort and her own innate ability, her colleague's "help" effort *and* ability, and a transitory shock. Thus, a worker's ability influences her colleague's performance. Workers consider work and help provision as two separate tasks and their cost-of-effort functions are task-specific. Workers' abilities and the shocks in production are independent and normally distributed. We also consider different degrees of incoming and outgoing skills benefit and peer interactions to and from a worker, which indicate how sensitive a worker's output is to her colleague's ability and help effort. All these parameters are exogenous.

First, we assume that multiperiod contracts are feasible, implying that the manager can commit herself to a life-time income path before the realization of production. The optimal contractual parameter based on a worker's own project output is always positive, indicating that a higher performance is

³Sabotage has been documented in several work environments. For example, political sabotage from party-mates, not necessarily from opponents, has shaped political campaigns and elections. Sabotage also happens in research labs. On April 26, 2011, the Office of Research Integrity issued a finding of research misconduct for Vipul Bhrigu, a former postdoctoral fellow at the University of Michigan. Bhrigu was caught on videotape sabotaging the experiments of Heather Ames, a graduate student in *his* lab. Heather was doing basic cancer research and began noticing problems with her research materials: errant antibodies dumped into her western blots, and several instances of ethanol in her cell culture media. Bhrigu was caught to intentionally undermine her work.

⁴With technological advances, sabotage is becoming easier to prove in some respects. Most computers store what is called "metadata", which includes information about computer-generated documents such as the date and time it was drafted or modified. In modern lawsuits, it is possible to discover metadata to cover up the truth.

⁵The existing literature uses such contracts when the market shocks that hit workers' production are correlated. In our setting, market shocks are independent. However, individual outputs are correlated due to colleague interactions. Relative performance evaluation schemes are the consequence of the efficient use of information conveyed by both performance measures about a worker's effort and ability.

rewarded with a higher payment.⁶ However, the sign of the optimal pay-for-peer performance parameter is not straightforward. We argue that the manager will induce low risk-averse workers to help each other, inducing cooperation between them, but this may not be the case for a highly risk-averse worker whose contribution in her colleague's production is small. In our model, both performance measures are sensitive to both worker's unknown abilities, increasing the variance of the rewards. Due to risk-sharing, the manager needs to provide insurance that involves a negative help effort, inducing sabotage. It aims to decrease the variance of a worker's wage and thus the required insurance. Because of sabotage, both workers also have stronger incentives to focus on their own projects and improve their own production. Thus, the manager is benefited because the cost of exerting effort decreases while the total production increases. This is in sharp contrast to Auriol et al. (2002) who assume that a colleagues' support depends only on her effort (not her ability). They argue that the manager, who provides long-term contracts, always induces the workers to help each other, regardless of their degree of tolerance against risk.

Second, we establish that under full commitment, while a manager can induce a worker to help her colleague when tenure is short, she may induce the same worker to sabotage if tenure is long. That is because the covariance of a worker's wages that are provided in different periods depends on the variance of both workers' abilities. As employment extends to many periods, the risk associated with both workers' abilities increases and influences significantly the manager's perception about workers' effort levels. Thus, a manager who establishes cooperation between the co-workers when employment is short, she may allow for sabotage as the duration of long-term contracts increases.

We analyze the conditions under which incentives to help or sabotage arise, attempting to better understand team formation. If the workers have incentives to help each other, the manager has incentives to keep the team together and commit to a multi-period contract. Help arises when workers contribute to peer production much with their effort and little with their ability. Thus, either a worker's ability should not be a key determinant of her colleague's output or the variance of abilities should be small enough. For example, the market has more information about senior workers, rather than juniors. The manager will also prefer to appoint workers with work experience suitable to both projects in order the effect of their help effort to be positive and significant. Thus, the manager can hire senior workers for whom there is information about their skills and also have similar work experience.

Third, assuming that commitment is not feasible, the manager offers short-term contracts that give rise to career concerns. She now renegotiates the contracts with the workers in each period and the market draws inferences about abilities via both workers' past project outputs. By exerting effort, a worker can influence her own and her colleague's performance measures in order to bias the learning process in her favor. A worker now wants to help or sabotage her colleague in her attempt to induce

⁶Some principal-agent models allow both parties to hold some bargaining power (e.g., Pitchford (1998)) while other models assume that either party can make a 'take-it-or-leave-it' offer (e.g., Mookherjee & Ray (2002)).

an upward revision of the estimate of her own ability and thus increase her future remuneration. She wants to look as productive as possible in absolute and relative terms. The manager now uses explicit (short-term) contracts to weaken a worker's reputational incentives to help or her desire to sabotage.

Different compensation contracts are offered to temporary workers who will stay in the firm but be paired with another worker in the next period or be hired by a different firm. Temporary workers know that they are unable to capitalize any change in market beliefs about her current colleague's ability. However, in our model, reputational incentives to help or sabotage also arise for a temporary worker because her own reputation depends on her *current* colleague's output. In Auriol et al. (2002), any implicit incentives to influence a teammate's production disappears in the case of temporary workers.

This paper is tied to the literature on career concerns based on Holmström (1982a). In his single-agent model, career concerns motivate a worker who has bargaining power vis-à-vis the market to work harder in the current period in order to build up her reputation, seeking for higher future rewards. In a two-agent model, Meyer & Vickers (1997) show that a worker wants to increase her reputational bonus but a ratchet effect arises. This happens because in her attempt to convince the market that she is of higher ability, the worker increases the expectations with respect to her production, making the manager to become more demanding. Thus, assuming that workers' abilities are positively correlated, each worker free-rides on the effort of the other to enhance reputation. Auriol et al. (2002) use relative performance evaluations, and consider work and help as two separate tasks, while a worker's output depends *only* on her *own* ability. The process of inference of each worker's ability is independent. They examine 'passive' sabotage because agents become reluctant to help their colleagues. Lazear (1989) considers sabotage incentives in tournaments.⁷ In our setting, the workers' innate characteristics are independent but they enter in both workers' production functions. We show that incentives to help or sabotage arise through career concerns and explicit contracts to both long-term and temporary workers. Nevertheless, the channels through which explicit incentives are determined in these two contractual environments are different.

We also contribute to the existing literature on moral hazard in teams.⁸ In a multi-agent environment where individual outputs are observable and correlated, a worker's compensation contract is made contingent on her colleague's production. Holmström (1982b) and Mookherjee (1984) show that rela-

⁷Bernhardt (1995) studies how the unobservability and composition of a worker's skills affect wage and promotion paths. Bar-Isaac & Deb (2014) depart from the assumption of homogeneous 'audiences' who will assess workers' abilities and examine career concerns when the audiences have diverse preferences. Ferrer (2010) studies the effects of lawyers' career concerns on litigation when the outcome of a trial depends on the opposing lawyers' effort and abilities. Cisternas (2017) considers that a worker's skills follow a Gaussian process with an endogenous component reflecting human-capital accumulation and discusses the worker's decisions to exert effort or invest in skill acquisition. Bilanakos (2013) argues that the provision of general training increases the worker's bargaining power vis-à-vis the employer.

⁸This paper is related to the literature on team incentives that examines the degree of visibility of workers' characteristics; i.e. Ortega (2003), Jeon (1996), Bar-Isaac (2007), Arya & Mittendorf (2011), Dewatripont, Jewitt & Tirole (1999), Dewatripont, Jewitt & Tirole (2000), Effinger & Polborn (2001), Milgrom & Oster (1987), Mukherjee (2008).

tive performance evaluation (RPE) schemes can filter out the common shock from a worker's reward, assuming that there are no technological interactions among workers. Since the workers are exposed to lower risk, the trade-off between insurance and incentives is shifted, facilitating stronger explicit incentive schemes. Itoh (1991, 1992) shows that a manager who uses RPE can install cooperation among workers only if the correlation between the workers' performances is low. We assume that the abilities and production shocks are uncorrelated but individual productions depend on both workers' innate abilities which are unknown, as in Chalioti (2016). We investigate how effectively a manager can induce cooperation or competition between workers through RPE when principal's commitment power changes.

This paper also complements the literature on the substitutive relationship between explicit and implicit incentives based on Gibbons & Murphy (1992), due to the additive production technology. Dewatripont et al. (1999) show that explicit and implicit (reputational) incentives may become complements, assuming that ability and effort enter the production function in a multiplicative fashion. Meyer, Olsen & Torsvik (1996) state that the ratchet effect can be weakened by inducing teamwork. In our setting where RPE schemes are provided and a worker's ability influences both workers' performance measures, a worker's implicit incentives to help or sabotage depend on both contractual parameters of her future explicit compensation.

The paper is organized as follows. Section 2 describes the model and discusses the production technologies. Section 3 analyzes the contracts when the manager can fully commit herself to a life-time income path. It analyzes the optimal explicit incentives in two-period and multi-period settings. Section 4 studies short-term contracts when workers stay together as long as employment lasts. It discusses the process of inference of workers' abilities and the trade-off between explicit and implicit incentives for these long-term workers. We also focus on the explicit and implicit motivation for temporary workers. Section 5 concludes.

2 The model

There are two risk averse workers, denoted by i and j , where $i \neq j$. They live for two periods indexed by $t = \{1, 2\}$ and they do not discount the future. They are also appointed by a risk neutral and profit seeking manager who compensates them with explicit contracts. At each period t , each individual works to fulfill her own project, while interacting with her colleague. Both project outputs are observable and contractible, allowing the manager to deal with each worker separately as in Itoh (1991, 1992).

2.1 Production technology

In period t , each worker i controls a stochastic production process. Her project output, z_t^i , is the sum of her own innate ability, θ^i , her nonnegative *work* effort, e_t^i , her colleague's support and a transitory shock, ε_t^i . Worker j 's support depends on her own ability, θ^j , weighted by a parameter $h_j \in [0, 1]$, and her *help* effort, a_t^j , weighted by $k_j \in [0, 1]$. Thus, worker i 's production function is⁹

$$z_t^i = \theta^i + e_t^i + \varepsilon_t^i + h_j \theta^j + k_j a_t^j. \quad (1)$$

Before production takes place, all parties have symmetric but imperfect information about workers' abilities as in Holmström (1982). However, the workers and all prospective employers believe that θ^i is drawn from a normal distribution with zero mean and variance σ_i^2 , where $\sigma_i^2 < \sigma_j^2$. Thus, there is less uncertainty about worker i 's attribute. The abilities are also independent of each other and of the noise terms. Worker i 's attribute θ^i can manifest the worker's ability to successfully accomplish a project which is symmetrically unknown to all parties at each stage.¹⁰ The random terms ε_t^i , ε_t^j follow a normal distribution with zero mean and variance φ_i^2 and φ_j^2 , respectively, where $\varphi_i^2 < \varphi_j^2$. They are also independently distributed across workers and periods.

In this model, a worker's ability and "help" effort enter her colleague's production process additively. Note that although a worker always exerts (positive) work effort to improve her own project outcome, her career concerns and the principal-agent problems may induce her, in equilibrium, either to help or even sabotage her colleague. Thus, workers' help efforts a_t^i and a_t^j can be negative. The parameter h_j measures the degree of *incoming skills benefit* to worker i that is generated by working together with a worker of ability θ^j . It indicates how sensitive worker i 's output is to her colleague's attribute. Similarly, h_i denotes the degree of *outgoing skills benefit* from worker i to worker j . These parameters are exogenous and less than one, so are k_i and k_j . The parameter k_j captures the degree of *incoming peer interactions* which is the fraction of worker j 's help effort that enters worker i 's project output, while k_i represents the degree of *outgoing peer interactions*.

The nature of interactions is imperfect, implying that putting effort into a worker's own task is more productive than providing help to a fellow worker. The level of all parameters affects each worker's ability to affect her colleague's production as well as to manipulate market perceptions about abilities. Note that when $k_i = 0$ and $h_i > 0$, a worker i does not actively influence her colleague's production but involuntarily she benefits it through her ability. By having the chance to observe the certain ways a colleague runs experiments in a lab, hearing the comments she makes in a seminar on someone else's

⁹We normalize the price of each worker's output to one and in all periods, we assume that the scale of production remains the same.

¹⁰Some papers study the optimal explicit incentives when the worker is better informed about her own characteristics than the market (Laffont & Tirole (1988)).

work or the ideas she shares in a meeting about a business project, a worker is benefited nonetheless.

2.2 Workers' preferences and objectives

Worker i receives the reward w_t^i and has constant absolute risk-averse (CARA) preferences. She is endowed with the utility function

$$U^i = -\exp \left\{ -r_i \sum_{t=1}^2 [w_t^i - g(e_t^i) - y(a_t^i)] \right\}, \quad (2)$$

where r_i is the Arrow-Pratt measure of risk aversion, $0 < r_i < r_j$. Due to the additive separability of the utility function, workers do not consider income smoothing across periods. They make their decisions as if they enjoy access to perfect capital markets in each period.

The contracts depend linearly on both workers' project outputs since the performance measures are correlated due to peer interactions. Holmström & Milgrom (1987) establish that in a model much like the single-period version of this model (but lacking the uncertainty about a worker's ability), the optimal contract is linear.¹¹ Gibbons & Murphy (1992), in a single-agent model, and Auriol et al. (2002), in their multi-agent framework, also consider contracts that are linear in outputs. At each period t , the manager offers contracts of the form $C_t^i \equiv (\omega_t^i, \beta_t^i, \gamma_t^i)$ and worker i receives

$$w_t^i = \omega_t^i + \beta_t^i z_t^i + \gamma_t^i z_t^j, \quad (3)$$

where ω_t^i denotes the fixed salary component and β_t^i, γ_t^i are the incentive parameters. Such incentive schemes introduce either cooperation or competition between the workers, depending on the sign of γ_t^i .

Risk-aversion on the part of the workers is essential when explicit contracts are provided so that the incentive parameters are less than one. Otherwise, the optimal contract will impose substantial human capital risk on the workers. Besides, because of workers' risk aversion, optimal contracts do not completely eliminate career concerns. As explicit payments decrease due to risk-sharing, reputational incentives increase.

We assume that workers' disutilities are task specific, as in the multi-agent models of Auriol et al. (2002), and Itoh (1992). The cost functions of work effort and help effort, $g(e_t^i)$ and $y(a_t^i)$ respectively, are twice continuously differentiable and convex, implying that there are diminishing returns to scale in the production process. A worker can also find less costly to exert effort to influence her colleague's production than putting effort to accomplish her own project. The cost of transmitting false information in order to confuse or distract a colleague can be very small. With task-specific cost functions, a

¹¹Holmström & Milgrom (1987) show that in a static version of their dynamic model, the optimal compensation scheme that is offered to a worker with CARA preferences is a linear function of the performance measures.

worker can focus on eliciting effort to affect her colleague's project output without having to consider simultaneously technologically founded externalities. Putting effort in a task does not require effort away from the other task. We also assume that $g'(0) = 0$, $y'(0) = 0$, $\lim_{e_t^i \rightarrow \infty} g'(e_t^i) = \infty$ and $\lim_{a_t^i \rightarrow \pm\infty} y'(a_t^i) = \pm\infty$.¹²

Under full information, the manager observes the workers' effort levels and thus can make contract offers that achieve specific effort assignments. In particular, worker i 's payment is fixed and equal to the sum of the cost of both efforts, $g(e_t^i) + y(a_t^i)$. The first-best levels of work and help efforts, $e_t^{i,FB}$ and $a_t^{i,FB}$, satisfy the conditions $g'(e_t^{i,FB}) = 1$ and $y'(a_t^{i,FB}) = k_i$, respectively. The workers receive the same contract in each period which is the repetition of the optimal contract in a one-shot game.

3 Long-term contracts

In section 3, we assume that long-term contracts are feasible. In a two-period model, the manager can commit herself to a second-period salary before the observation of the first-period outputs. Thus, the manager can insulate a worker's expected life-time compensation from the risk associated with true abilities - actual θ^i and θ^j . A worker's problem is identical across time and career concerns do not arise. Notice that a long-term contract is not identical with two one-period contracts.

3.1 Two-period model

The manager's net profit equals the sum of the project outputs net of workers' compensations. In a two-period model, the manager's problem is

$$\max_{C_t^i, e_t^i, a_t^i, C_t^j, e_t^j, a_t^j} E \{U^P\} = E \left\{ \sum_{t=1}^2 \sum_{i=1}^2 (z_t^i - w_t^i) \right\}$$

$$\text{subject to } CE^i \equiv \sum_{t=1}^2 [E \{w_t^i\} - g(e_t^i) - y(a_t^i) - \frac{r_i}{2} \text{Var} \{w_1^i + w_2^i\}] \geq 0, \quad \forall i \quad (IR^i)$$

$$e_t^{i*} = \arg \max_{e_t^i} CE^i, \quad \forall i, t \quad (IC_{e,t}^i)$$

$$a_t^{i*} = \arg \max_{a_t^i} CE^i, \quad \forall i, t \quad (IC_{a,t}^i)$$

$$\text{where } \text{Var} \{w_1^i + w_2^i\} = \left[\sum_{t=1}^2 (\beta_t^i + h_i \gamma_t^i) \right]^2 \sigma_i^2 + \left[\sum_{t=1}^2 (h_j \beta_t^i + \gamma_t^i) \right]^2 \sigma_j^2 + \sum_{t=1}^2 [(\beta_t^i)^2 \varphi_i^2 + (\gamma_t^i)^2 \varphi_j^2].$$

¹²Models based on Holmström & Milgrom (1991) assume total-effort-cost functions. The cross-partial derivatives with respect to two tasks are positive. As a worker increases the effort directed to one task, the marginal cost of effort to the other task will also increase. Thus, exerting help effort would be costly to a worker and it crowds out effort directed to her own project, decreasing her own production. In equilibrium, each worker equates the marginal return to effort in both tasks.

The individual rationality (IR^i) constraint shows that the worker will participate in the production process only if the certainty equivalence of her utility, CE^i , exceeds her outside option.¹³ Since the manager specifies the long-term incentive schemes before the realization of the first period outputs, the outside option is equal to her expected innate ability, which is normalized to zero.

In this framework where workers' utility is additively separable across time, and current production is independent of the workers' past outcomes, the manager cannot exploit any gains from intertemporal risk sharing. Instead, the second-period rewards depend only on the second-period outcomes, ignoring the first-period production. Thus, the same payment schemes are provided in both periods. The IR^i constraint is binding at the optimum. Notice that only the sum of the base payments in both periods, $\omega_{1,f}^i + \omega_{2,f}^i$, is determined by the optimal contract; the subscript f denotes the equilibrium values in the full commitment case. Because of this feature, the manager can guarantee that the workers will not exercise their right to quit before the end of this relationship. By setting a very low first-period base payment and a very large second-period base payment, the manager can costlessly make unprofitable the option of quitting for both workers.

The incentive compatibility constraints, $IC_{e,t}^i$ and $IC_{a,t}^i$, guarantee that a worker chooses the (expected) utility maximizing efforts. The optimal work and help efforts satisfy, respectively,¹⁴

$$\beta_t^i = g'(e_t^i) \text{ and } k_i \gamma_t^i = y'(a_t^i), \quad \forall i. \quad (4)$$

The long-term explicit incentives are¹⁵

$$\beta_f^i = \frac{1}{1 + r_i (\Sigma_f^{ii} + \Phi_f^i \Sigma_f^{ij}) g_f^{i''}} \text{ and } \gamma_f^i = \Phi_f^i \beta_f^i, \quad (5)$$

where $\Phi_f^i \equiv \frac{k_i^2(1+r_i \Sigma_f^{ii} g_f^{i''}) - 2r_i \Sigma_f^{ij} y_f^{j''}}{k_i^2(1-2r_i \Sigma_f^{ij} g_f^{i''}) + r_i \Sigma_f^{jj} y_f^{j''}}$, $\Sigma_f^{ii} \equiv \varphi_i^2 + 2(\sigma_i^2 + h_j^2 \sigma_j^2)$ and $\Sigma_f^{ij} \equiv h_i \sigma_i^2 + h_j \sigma_j^2$. The second derivatives $g_f^{i''}$ and $y_f^{j''}$ are always positive because of the convexity of the cost-of-effort functions. Proposition 1 states the conditions under which the manager will find it optimal to induce a worker to help or sabotage her colleague. This is in sharp contrast to Auriol et al. (2002) where the long-term explicit incentives based on both workers' performance measures are always positive, inducing a worker always to help her colleague.

¹³In this multi-agent framework, the monotone likelihood ratio property (MLRP) and the convexity of the distribution function condition (CDFC) are not sufficient for the first-order approach to be valid as in a single-agent setting. Itoh (1991) argues that, in a model with cross-agent interactions, a generalized CDFC for the joint probability distribution of the outputs is needed and the wage schemes must be nondecreasing. The coefficient of absolute risk aversion must also not decline too fast. Our model with workers' CARA preferences, linear contracts and production technologies satisfies all these assumptions and thus the first-order approach applies.

¹⁴The manager's problem is solved in appendix A.1.

¹⁵In case of perfect peer interactions and identical skills benefit, $h_i = h_j = k_i = k_j = 1$, where there are no frictions in interacting with a coworker, both incentive parameters in (5) are identical. Wages are contingent on the total output and only the observation of the aggregate measure $z_t^i + z_t^j$ is needed.

Proposition 1 (Long-term explicit incentives) *Given that the manager fully commits to a worker's life-time income path,*

(a) *the equilibrium pay-for-own-performance parameter is always positive, $\beta_f^i > 0$,*
 while (b) *the pay-for-peer-performance parameter is negative, $\gamma_f^i < 0$, inducing a worker to sabotage,*
if and only if $\frac{k_i^2(1+r_i\Sigma_f^{ii}g_f^{ii})}{2r_i y_f^{ii}} - h_j \frac{\sigma_j^2}{\sigma_i^2} < h_i$.

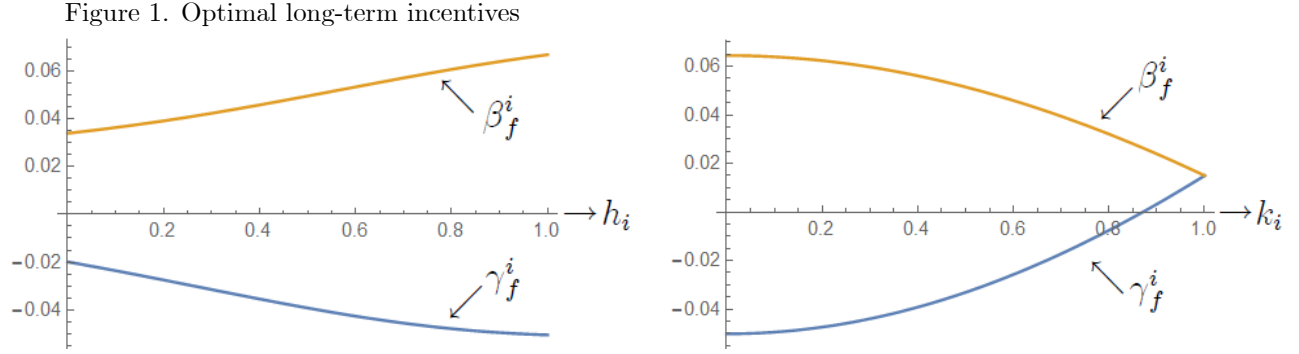
Proof. In Appendix (A.1). ■

The positive sign of β_f^i indicates that a worker's higher own project output is compensated with a higher wage. Due to risk-sharing, β_f^i falls short from its efficient level, so does γ_f^i . However, the sign of γ_f^i is less straightforward. The degree of peer interactions and the contribution of workers' abilities in performance measures play a key role in specifying the optimal long term explicit incentives.

The manager sets γ_f^i positive if worker j 's performance, z_t^j , is more sensitive to worker i 's help effort rather than on her ability (h_i lower than k_i). A risk-averse worker will require insurance against low realization of both workers' (unknown) abilities which affect both performance measures. On the one hand, a positive γ_f^i increases the variance of worker i 's wage, $Var\{w_1^i + w_2^i\}$, and thus the risk to which she is exposed. However, on the other hand, it increases the performance of her colleague, so does her wages, w_1^i and w_2^i . The optimal γ_f^i will be positive, incentivizing a worker to help, when the increase in the cost of exerting help effort is smaller than the increase in a colleague's output. Thus, the manager anticipates the support a worker provides to her colleague and rewards her when her colleague does better. Relative performance evaluation schemes can effectively be used as means of internalizing the positive effects of providing support.

The "compensation ratio" $\left| \frac{\gamma_f^i}{\beta_f^i} \right|$ becomes larger in compensation packages with a higher h_i and a lower k_i . As a worker's help effort weights more in her colleague's performance than the risk added because of her unknown ability, she is motivated to pay more attention to her colleague's project and improve her output. The manager is benefited by shifting worker i 's focus on improving her colleague's production. Besides, as the degree of the incoming skills benefit from worker j to worker i , measured by h_j , increases, β_f^i will decrease. Worker i 's output as a performance measure becomes noisier, implying that it conveys less information about worker i 's work effort. In turn, the manager relies less on z_t^i to anticipate e_t^i . Given also that the variance of her compensation increases due to higher risk which is introduced by a factor incorporated in z_t^i , a lower β_f^i will mitigate this effect. Thus, the manager considers to balance between the optimal incentives and the focus of the agent. Any change in k_j leaves worker i 's incentives unaffected since the channel through which her efforts influence the total production and the risk to which she is exposed only depend on h_i , h_j and k_i . The optimal γ_f^i , when positive, also increases with the uncertainty about worker i 's ability and the variance of worker i 's transitory shock, σ_i^2 and φ_i^2 , while

decreases with σ_j^2 and φ_j^2 .



Figures 1.a and 1.b show how the optimal long-term incentives change with h_i and k_i , respectively, when $\sigma_i^2 = \sigma_j^2 = \varphi_i^2 = \varphi_j^2 = 2$, $r_i = 5$, and $g_f''' = y_f''' = 1$. We assume $k_i = 0.9$ in Figure 1.a, while $h_i = 0.1$ in Figure 1.b.

For a high risk averse worker i whose help effort weights less than her ability in her colleague's performance (high h_i and low k_i), the manager sets γ_f^i negative. Negative explicit incentives provided by long-term contracts are in contrast to Auriol et al. (2002), in which the pay-for-peer-performance parameters are always positive. In our framework, such incentives can be reversed. The sensitivity of both performance measures on both colleagues' abilities, which are unknown, increases the variance of the rewards, and by setting γ_f^i negative, the manager aims exactly in decreasing the variance of worker i 's wage. Notice that because of a negative γ_f^i , worker i is induced to sabotage, while she is rewarded for the effort she put in sabotaging her colleague; i.e., worker i 's wage depends on $k_i \gamma_f^i a_t^i$ which is positive.

Given that the manager's benefit equals the sum of both projects' outputs, $z_t^i + z_t^j$ at each period t , one could consider that the manager would prefer to shut down any incentives of worker i to affect her colleague's output by setting γ_f^i zero. In this case, β_f^i is smaller, so is the total production $z_t^i + z_t^j$. It is optimal for a profit-seeking manager to set a negative γ_f^i and increase β_f^i so as to motivate worker i to focus on her own project. Because k_i is small, the effect of sabotage will also be small. Thus, we argue that the manager may allow for little sabotage in order to decrease the provided insurance, implying lower cost of exerting effort, and encourage a worker to focus on her own projects, increasing her own output. Even if the degree of peer interactions are equal, the intensity of optimal incentives and thus workers' focus of attention will depend on the amount of noise about each worker's abilities which is reflected by the variance of θ s.

Proposition 2 (Precision about abilities & optimal incentives) Suppose $k_i = k_j = h_i = h_j$, $r_i = r_j$, and that the manager has a better precision about worker i 's ability, $\sigma_i^2 < \sigma_j^2$. Then, the manager will offer an incentive scheme to worker i which is more individually oriented than the one given to worker j : $\beta_f^i > \beta_f^j$ and $|\gamma_f^i| < |\gamma_f^j|$.

Proposition 2 highlights that the manager will exert a higher work effort and a lower help effort by the worker with the lower variance of her ability. Since the manager has more information about θ^i , worker i is induced to focus more on the own project than influencing her colleague's production.

3.2 Multiperiod model

We now examine the effects of full commitment to a life-time salary path when employment extends to many periods. Suppose that the workers are appointed for τ periods, where $\tau > 2$. In this framework, the variance of the life-time payments is

$$Var \left\{ \sum_{t=1}^{\tau} w_t^i \right\} = \left[\sum_{t=1}^{\tau} (\beta_t^i + h_i \gamma_t^i) \right]^2 \sigma_i^2 + \left[\sum_{t=1}^{\tau} (h_j \beta_t^i + \gamma_t^i) \right]^2 \sigma_j^2 + \sum_{t=1}^{\tau} \left[(\beta_t^i)^2 \varphi_i^2 + (\gamma_t^i)^2 \varphi_j^2 \right].$$

In equilibrium, the contractual parameters become

$$\beta_{\tau,f}^i = \frac{1}{1 + r_i (\Sigma_{\tau,f}^{ii} + \Phi_{\tau,f}^i \Sigma_{\tau,f}^{ij}) g_f^{i''}} \text{ and } \gamma_{\tau,f}^i = \Phi_{\tau,f}^i \beta_{\tau,f}^i,$$

where $\Phi_{\tau,f}^i \equiv \frac{k_i^2 (1 + r_i \Sigma_{\tau,f}^{ii} g_{\tau,f}^{i''}) - \tau r_i \Sigma_{\tau,f}^{ij} y_{\tau,f}^{i''}}{k_i^2 (1 - \tau r_i \Sigma_{\tau,f}^{ij} g_{\tau,f}^{i''}) + r_i \Sigma_{\tau,f}^{jj} y_{\tau,f}^{i''}}$ and $\Sigma_{\tau,f}^{ii} \equiv \varphi_i^2 + \tau (\sigma_i^2 + h_j^2 \sigma_j^2)$. Proposition 3 highlights that for a short employment period, a worker can be induced to help her colleague, while for a long employment period, such incentives can be negative. Thus, sabotage incentives are more likely to prevail for highly risk-averse agents whose tenure is long.

Proposition 3 (Incentives & multiperiod employment) *Given that the manager fully commits to a life-time income path, equilibrium pay-for-peer-performance parameter is negative, $\gamma_{\tau,f}^i < 0$, if and only if worker i 's employment period is long enough,*

$$\frac{k_i^2 (1 + r_i \varphi_i^2 y_{\tau,f}^{i''})}{r_i \Sigma_{\tau,f}^{ij} y_{\tau,f}^{i''} - k_i^2 (\sigma_i^2 + h_j^2 \sigma_j^2) g_{\tau,f}^{i''}} < \tau.$$

Note that under the assumption of independently distributed random terms, the covariance of wages offered in different periods depends solely on σ_i^2 and σ_j^2 (not φ_i^2 and φ_j^2); i.e., for any t , $cov(w_t^i, w_{t+1}^i) = 2 [(\beta_t^i + h_i \gamma_t^i) (\beta_{t+1}^i + h_i \gamma_{t+1}^i) \sigma_i^2 + (h_j \beta_t^i + \gamma_t^i) (h_j \beta_{t+1}^i + \gamma_{t+1}^i) \sigma_j^2]$. Thus, for a longer employment period, the noise in the production processes introduced by a colleague's ability matters more in shaping workers' contracts. It affects a worker's production in many periods, increasing the risk to which the worker is exposed. More precisely, higher risk in longer-term contracts is compensated with weaker pay-for-own-performance incentives; i.e., $\frac{\partial \beta_{\tau,f}^i}{\partial \tau} < 0$ for all τ . Besides, although for a short employment period, the manager can incentivize a worker to help her colleague (positive $\gamma_{\tau,f}^i$), for a longer employment period, the manager will induce the same worker to sabotage (negative $\gamma_{\tau,f}^i$), seeking to

decrease the variance of her wage and make her focus more on her own project. Thus, contracts with different duration can provide opposing incentives to the same worker on how to influence a colleague's production.

This analysis indicates that the manager would prefer to hire low-risk averse workers with work experience suitable to both projects, so that their efforts have a significant impact in both project outputs (high k s). The variance of their θ s should also be small enough. Having formed a pair whose members have incentives to help each other, the manager will prefer to sign long-term contracts with them for many years of employment - i.e., manager's net payoffs increase with τ .¹⁶

4 Short-term contracts

We now examine the explicit incentives in a 2-period setting where the manager cannot commit herself to a life-time compensation scheme. Instead, she renegotiates a contract offer in each period. Career concerns arise and substitute explicit motivation.

4.1 Timing of the game and reputational bonus

The timing of the game has as follows. In the beginning of period 1, the manager simultaneously makes a contract offer to each worker. If a worker accepts the offer, she makes the effort choices. Events beyond the workers' control occur, both project outputs are realized and the contracts are executed. In period 2, all parties observe the realization of z_1^i and z_1^j , and update their assessments about abilities. Thus, past production affects future remuneration and incentives. A new contract offer is made to each worker and if she stays in the firm, she chooses the new effort levels. After the observation of the current production, the rewards are paid.

If a worker rejects the contract offer, she receives her outside option that equals her reputational bonus

$$\Theta_t^i \equiv (1 + h_i) E \{ \theta^i \mid z_{t-1}^i, z_{t-1}^j \} + \widehat{e}_t^i + k_i \widehat{a}_t^i.$$

This is a fixed payment which equals the total rents each worker can claim for her contribution to both workers' project outputs. Given the available information, her payment increases with an upward revision of the market's estimate of her own ability.

4.2 Learning process and career concerns

All parties observe past production in order to infer the level of workers' abilities. When output shocks are uncorrelated and a worker's ability does not affect her colleague's performance as in Auriol

¹⁶For example, Hahn (2017) recently studied committee design in a different setup.

et al. (2002), the process of inference of each worker's ability is independent. However, in this model, as in Chalioti (2016), both performance measures, z_1^i and z_1^j , convey information about θ^i and thus are used in the updating process about its actual level. The conditional distribution of θ^i is normal with mean

$$E \{ \theta^i \mid z_1^i, z_1^j \} = \rho_1^{ii} (z_1^i - \widehat{e}_1^i - k_j \widehat{a}_1^j) + \rho_1^{ij} (z_1^j - \widehat{e}_1^j - k_i \widehat{a}_1^i),$$

and variance,

$$\sigma_{i,2}^2 \equiv \sigma_i^2 (1 - \rho_1^{ii} - h_i \rho_1^{ij}),$$

where $\widehat{e}_1^i$ and $\widehat{a}_1^i$ are the conjectures of worker i 's current efforts.¹⁷ In comparison to the full commitment case, the variance of each worker's second-period production is smaller because all parties observe past performance and thus have more precise predictions about workers' abilities. The conditional correlation coefficients are, respectively,

$$\begin{aligned} \rho_1^{ii} &\equiv \text{corr} \{ \theta^i, z_1^i \mid z_1^j \} = \frac{\sigma_i^2}{\lambda_1} [\varphi_i^2 + (1 - h_i h_j) \sigma_j^2], \\ \rho_1^{ij} &\equiv \text{corr} \{ \theta^i, z_1^j \mid z_1^i \} = \frac{\sigma_i^2}{\lambda_1} [h_i \varphi_j^2 - (1 - h_i h_j) h_j \sigma_j^2], \end{aligned}$$

where $\lambda_1 = \varphi_i^2 \varphi_j^2 + (1 - h_i h_j)^2 \sigma_i^2 \sigma_j^2 + \varphi_i^2 (\sigma_j^2 + h_i^2 \sigma_i^2) + \varphi_j^2 (\sigma_i^2 + h_j^2 \sigma_j^2)$.

The signal ρ_1^{ii} is always positive, $\rho_1^{ii} > 0$, implying that given the realization of a colleague's production, a worker's high own performance signals high own ability and vice versa. The signal ρ_1^{ij} is also positive as long as worker i 's ability affects significantly her colleague's performance, while her own performance is not very sensitive to her colleague's ability (high h_i and low h_j). In this case, given z_1^i , a higher colleague's output z_1^j is attributed to worker i 's ability. It is a good signal for her, resulting in an upward revision of the market estimate of θ^i . Effort is a substitute for ability, implying that a worker can manipulate market perceptions about abilities by distorting effort levels. Hence, in the absence of explicit motivation where a worker's reward exactly equals her reputational bonus, when ρ_1^{ij} is positive, worker i will have incentives to *help* her colleague in order to build up her reputation. Worker i 's help will increase z_1^j , shaping market assessments about her own ability in her favor and thus increasing worker's future remuneration.

The opposite occurs in a setting with a small outgoing skills benefit, h_i , but high h_j . Both signals are now sensitive to θ^j , while the impact of worker i 's ability, θ^i , on worker j 's production is less significant. In this case, it is more likely that both performance measures reflect the level of coworker ability θ^j . If both workers perform well, the market attributes these good outcomes to high θ^j , causing the estimate

¹⁷ Assuming also that all parties have rational expectations, we have $\widehat{e}_1^i = e_1^{i*}$ and $\widehat{a}_1^i = a_1^{i*}$. The equilibrium conjectures must be correct. The presence of noise insures that there are no off-equilibrium realizations of observables, and in equilibrium, each agent is restricted to exert the levels of effort that are expected of her. Underprovision of effort will influence the updating process against her.

of θ^i to be revised downwards. A higher output from worker j is now a bad signal of worker i 's ability. By helping a colleague to further increase her project output, worker i will induce market inferences to be revised against her. Instead, bad performance by her colleague will be a good signal of her own ability. A decrease in z_1^j will increase worker i 's reputation so that even when explicit contracts are not provided, she will now have incentives to sabotage her colleague because of career concerns. She wants to sabotage in her attempts to manipulate the learning process in her favor.

4.3 Permanent workers

The manager's problem is the same as in the full commitment case, but with a different individual rationality (IR_t^i) constraint in each period t ,

$$CE_t^i \equiv \sum_{t=1}^2 \left[E \{ w_t^i \mid z_{t-1}^i, z_{t-1}^j \} - g(e_t^i) - y(a_t^i) - \frac{r_i}{2} Var \left\{ \sum_{t=1}^2 w_t^i \mid z_{t-1}^i, z_{t-1}^j \right\} \right] \geq \Theta_t^i.$$

In period 2, each worker i can now take a report of both her own and her colleague's first-period outputs to her manager and other prospective employers in order to let them compare the performances. Given both reports, the manager offers to each worker i her reputational bonus $\Theta_2^i = (1 + h_i) E \{ \theta^i \mid z_1^i, z_1^j \} + \widehat{e}_2^i + k_i \widehat{a}_2^i$. The IR_2^i constraint is binding at the optimum, implying that the manager is equally well-off by hiring either a high reputation worker at a high wage or a low reputation worker at a low wage. Each worker knows that competing employers cannot make a better offer than Θ_2^i .¹⁸ This bargaining outcome can arise as the equilibrium of an extensive-form game. In this game, a worker is randomly assigned to a prospective employer and queues with the other job applicants. The employer makes a contract offer to the first worker in line. If the worker accepts the offer, she works for this employer. Otherwise, she queues for another job and the employer makes an offer to the next worker in line. Therefore, each worker receives only her own reputational bonus. The manager signs the most appealing contracts.

The second period in the short-term contracting framework is isomorphic to the incentives raised in the first period in the full-commitment case, analyzed in section 3. In particular, the optimal work and help efforts, e_2^{i*} and a_2^{i*} , satisfy equations (4). We also compute the base payment by solving the IR_2^i constraint when binding:

$$\omega_2^i = \Theta_2^i - E \{ \beta_2^i z_2^i + \gamma_2^i z_2^j \mid z_1^i, z_1^j \} + g(e_2^{i*}) + y(a_2^{i*}) + \frac{r_i}{2} Var \{ w_2^i \mid z_1^i, z_1^j \}.$$

¹⁸Gibbons & Murphy (1992), Holmström (1999), among others, assume in single-agent models that the workers have all the bargaining power, and thus the manager maximizes her payoffs subject to a zero-profit condition. In a multi-agent setting, following Auriol et al. (2002), we consider a bargaining process that effectively makes each worker the residual claimant only to her reputational bonus. Nevertheless, the results would qualitatively be the same in both bargaining environments.

The optimal contractual parameters β_2^{i*} and γ_2^{i*} can be obtained by (5) replacing Σ^{ii} with $\Sigma_2^{ii} \equiv \varphi_i^2 + h_i^2 \sigma_{i,2}^2 + \sigma_{j,2}^2$ and Σ^{ij} with $\Sigma_2^{ij} \equiv h_i \sigma_{i,2}^2 + h_j \sigma_{j,2}^2$. We also have $\Phi_2^i \equiv \frac{k_i^2(1+r_i \Sigma_2^{ii} g_2^{i''}) - r_i \Sigma_2^{ij} y_2^{j''}}{k_i^2(1-r_i \Sigma_2^{ij} g_2^{i''}) + r_i \Sigma_2^{jj} y_2^{j''}}$.

In period 1, workers have explicit incentives from the compensation contracts *and* implicit incentives from career concerns. There is an implicit dependence of the future wage on the current project outputs. Current efforts affect the intercept of future wage, ω_2^i , but not the incentive parameters β_2^{i*} and γ_2^{i*} because there are no wealth effects in worker utility and the production functions are additive. Both workers have the same marginal product of effort regardless of their true ability. The optimal explicit incentives depend on the variance of outputs but not on their mean. As a result, worker i chooses the effort levels that satisfy the equations

$$\beta_1^i + M_P^{ii} = g'(e_1^{i*}) \text{ and } k_i(\gamma_1^i + M_P^{ij}) = y'(a_1^{i*}),$$

where $\frac{\partial \omega_2^{i*}}{\partial e_1^i} \equiv M_P^{ii}$ and $\frac{\partial \omega_2^{i*}}{\partial a_1^i} \equiv k_i M_P^{ij}$. A permanent worker i 's implicit incentives that arise through work and help efforts are, respectively,

$$M_P^{ii} = (1 + h_i - \beta_2^{i*} - h_i \gamma_2^{i*}) \rho_1^{ii} - (h_j \beta_2^{i*} + \gamma_2^{i*}) \rho_1^{ji}, \quad (6)$$

$$M_P^{ij} = (1 + h_i - \beta_2^{i*} - h_i \gamma_2^{i*}) \rho_1^{ij} - (h_j \beta_2^{i*} + \gamma_2^{i*}) \rho_1^{jj}. \quad (7)$$

Notice that the terms $(\beta_2^{i*} + h_i \gamma_2^{i*}) \rho_1^{ii}$ and $(h_j \beta_2^{i*} + \gamma_2^{i*}) \rho_1^{ji}$ in equation (6) arise due to the effect of e_1^i on $E\{\theta_2^i | z_1^i, z_1^j\}$ and $E\{\theta_2^j | z_1^i, z_1^j\}$ through z_1^i . Similarly, the terms $(\beta_2^{i*} + h_i \gamma_2^{i*}) \rho_1^{ij}$ and $(h_j \beta_2^{i*} + \gamma_2^{i*}) \rho_1^{jj}$ in equation (7) arise due to the effect of a_1^i on the conditional expectations of θ_2^i and θ_2^j through z_1^j .

The implicit incentives captured by M_P^{ii} indicate that by exerting more work effort in the first period, a worker aims to increase her reputational bonus by $(1 + h_i) \rho_1^{ii}$. However, this bonus is diminished by $(\beta_2^{i*} + h_i \gamma_2^{i*}) \rho_1^{ii} + (h_j \beta_2^{i*} + \gamma_2^{i*}) \rho_1^{ji}$. Suppose that the outgoing peer interactions from a low risk averse worker are large (high k_i) so that $\gamma_2^{i*} > 0$. If $\rho_1^{ji} \equiv \text{corr}\{\theta^j, z_1^i | z_1^j\}$ is positive, the manager infers that worker j has reputational incentives to help her co-worker i and increase her project output. Thus, because worker i is going to be assessed as being of higher ability and the explicit incentive component of her future remuneration is going to be large, the manager offers a contract whose base payment increases by less than the increase in the worker's reputational bonus. Notice though that if ρ_1^{ji} is negative, worker j has reputational incentives to sabotage. The manager now anticipates such incentives and adjusts the base payment upwards by $(h_j \beta_2^{i*} + \gamma_2^{i*}) \rho_1^{ji}$. Thus, the manager takes into account that sabotage by a colleague makes it harder for a worker to perform well and enjoy a high output. Besides, the base payment is also adjusted upwards for a high risk averse worker who receives a contract with $\gamma_2^{i*} < 0$. Note that $h_i \rho_1^{ii} + \rho_1^{ji}$ in equation (6) is positive. Hence, because of the trade-off between effort provision and insurance, the explicit component of the second-period contract is going to be smaller, resulting in

a higher base payment.

Worker i 's reputational incentives to influence her colleague's output are captured by M_P^{ij} . The most interesting cases are when $\rho_1^{ij} < 0$. In these cases, if a worker helps her colleague to increase the latter's output, she will make herself worse off, because it will shape market assessments against her. The manager will perceive that she is paired with a high productivity worker. Thus, it is in worker i 's interest to convince the manager the opposite and to do so, she sabotage her colleague. A negative γ_2^{i*} is going to make worker i less eager to sabotage in her attempt to build up her reputation. The manager can now use a negative γ_2^{i*} as a tool to decrease the risk to which the worker is exposed as well as reduce her appetite to sabotage because of her career concerns.

To derive the first period explicit incentives, we denote $B_1^i \equiv \beta_1^i + M_P^{ii}$ and $\Gamma_1^i \equiv \gamma_1^i + M_P^{ij}$. The equilibrium incentive parameters in the first period are

$$\beta_1^{i*} = \frac{1}{\Omega_1^i} - M_P^{ii} - r_i g_1^{i''} \frac{(k_i^2 + r_i \varphi_j^2 y_1^{i''}) [(\sigma_i^2 + h_j^2 \sigma_j^2) \beta_2^{i*} + \Sigma_1^{ij} \gamma_2^{i*}] + r_i (1 - h_i h_j)^2 \sigma_i^2 \sigma_j^2 g_1^{i''} \beta_2^{i*}}{\zeta_1^i \Omega_1^i}, \quad (8)$$

$$\gamma_1^{i*} = \frac{\Phi_1^i}{\Omega_1^i} - M_P^{ij} - r_i y_1^{i''} \frac{(1 + r_i \varphi_i^2 g_1^{i''}) [\Sigma_1^{ij} \beta_2^{i*} + (\sigma_j^2 + h_i^2 \sigma_i^2) \gamma_2^{i*}] + r_i (1 - h_i h_j)^2 \sigma_i^2 \sigma_j^2 g_1^{i''} \gamma_2^{i*}}{\zeta_1^i \Omega_1^i}, \quad (9)$$

where $\Omega_1^i \equiv 1 + r_i (\Sigma_1^{ii} + \Phi_1^i \Sigma_1^{ij}) g_1^{i''}$ and $\zeta_1^i \equiv k_i^2 (1 - r_i \Sigma_1^{ij} g_1^{i''}) + r_i \Sigma_1^{jj} y_1^{i''} > 0$ (see appendix (A.2)).

To better understand these results, we decompose the optimal explicit incentives and examine the underlying effects. The first term in (8) reflects a *noise reduction effect* that arises due to changes in the 'amount' of available information about ability. In the next period, as the market learns more about abilities and their conditional variances decrease, we have $\Omega_1^i > \Omega_2^i$. Therefore, learning reduces risk over time and stronger explicit incentives to work can be provided, $\beta_2^{i*} > \frac{1}{\Omega_1^i}$. Higher k_i shifts the incentive-insurance trade-off towards the former even more.

The last terms in both equations (8) and (9) capture *human capital insurance effects*: risk-aversion and uncertainty about abilities induce each worker to require insurance against low realizations of both θ^i and θ^j . The manager offers contracts that insure the workers against the intertemporal risk they face. When both β_2^{i*} and γ_2^{i*} are positive and large, a worker i incurs higher risk in the second period. Thus, the manager reduces the first-period incentives β_1^{i*} and γ_1^{i*} . If the manager uses a negative γ_2^{i*} as a tool to decrease the risk a worker i faces in the second-period, the decrease in β_1^{i*} and γ_1^{i*} that is caused by the human capital insurance effects is smaller.

The manager also adjusts the optimal explicit incentives to account for workers' career concerns. Equation (8) shows that the manager imposes a lower pay-for-own performance relation when the optimal reputational incentives to work are stronger. Similarly, equation (9) shows that when $M_P^{ij} > 0$, γ_1^{i*} is adjusted downwards. However, if a worker i has reputational incentives to sabotage her colleague, the optimal γ_1^i increases. Such incentives need to be undone by a higher γ_1^{i*} . The first term in equation

(9) captures the compensation ratio effect which shows how effective relative performance evaluation schemes are in inducing help or sabotage between the co-workers in the first period; i.e., if $M_P^{ii} = M_P^{ij} = \beta_2^{i*} = \gamma_2^{i*} = 0$, $\frac{\gamma_1^{i*}}{\beta_1^{i*}} = \Phi_1^i$.

The workers' optimal incentives are different from those in Lazear (1989) where sabotage incentives arise in tournaments because a worker's compensation is conditioned negatively to her colleagues' performances. Using such schemes, workers may desire to destroy other workers' output rather than to work harder on their own project. In our model, even if workers have incentives to sabotage because of career concerns, a negative γ_2^{i*} is used to decrease the risk to which a worker i is exposed and also to weaken worker's eagerness to sabotage. A negative γ_2^{i*} is also associated with a higher β_2^{i*} , motivating the worker to focus more on her own project. This result is also different from Meyer & Vickers (1997) in which the workers' abilities are correlated and reputational incentives are weakened due to the ratchet effect. In Auriol et al. (2002), reputational incentives to sabotage arise only when explicit contracts are provided, because workers become reluctant to help their colleagues. There is an implicit ratchet effect.¹⁹

4.4 Temporary workers

We have assumed so far that workers differ with respect to their innate abilities and they are assigned to work together as long as employment lasts. Tenure is long and a worker i sees her employment in a firm with a specific co-worker j as a life-time career. Instead, we now want to examine how reputational incentives and explicit contracting are shaped when employment or partnerships are temporary. In the end of each period, a worker can either enter the external labor market seeking for a new employer when her tenure in a specific firm is short, or stay in the firm but pair with a different worker of different ability. Workers do not now intent to manipulate the market assessments about her current colleague's ability since there are no benefits of convincing the market that she is paired with a low productivity worker. Market perceptions about θ^j play no role in worker i 's future remuneration. However, in our setting, reputational incentives to help or sabotage still arise because the *current colleague's output* is used in the updating process about a worker's *own* ability.

Each worker knows that, in each period, she will be paired with another worker or work for another firm. In Auriol et al. (2002), any reputational incentives to influence the current colleague's output disappear. Each worker knows that she is unable to capitalize any change in market beliefs about her

¹⁹Fama (1980)'s conclusion that explicit contracts are unnecessary to solve the principal-agent conflicts, since the market induces the "right" effort levels, also holds in our model. As in Holmström (1999), we can assume that $T \rightarrow \infty$ and workers discount the future with some factor $\delta \in [0, 1]$. Abilities fluctuate over time and remain unknown to the parties: $\theta_{t+1}^i = \theta_t^i + \eta_t^i$, where $\eta_t^i \sim N(0, \sigma_\eta^2)$ is independently distributed. In the supplementary material, we show that if there is no discounting, $\delta = 1$, and $r_i > 0$, the stationary explicit incentives are zero and efforts are efficient for any h_i , h_j , k_i and k_j . Bar-Isaac & Hörner (2014) and Bonatti & Hörner (2017) discuss the relationship between career concerns and market structure in different settings.

current colleague's ability while she has to bear the cost of exerting effort to bias these beliefs. The learning process about her own ability is independent from her current colleague's output. However, in our model, this is not the case.

Reputational incentives to work and help or sabotage still arise when the duration of pairing is short:

$$M_T^{ii} = (1 + h_i - \beta_2^{i*} - h_i \gamma_2^{i*}) \rho_1^{ii}, \quad (10)$$

$$M_T^{ij} = (1 + h_i - \beta_2^{i*} - h_i \gamma_2^{i*}) \rho_1^{ij}. \quad (11)$$

The comparison of equations (6) and (10) reveals that a temporary worker does not bother to influence her current colleague's expected ability, since she will not be paired with her in the following period. Thus, the last term in (6) which reflects the incentives of a long-term worker i to manipulate market perceptions about θ^j disappears. In our model, a worker who is appointed in a temporary basis still intends to exert α_1^i in order to affect her current colleague's output because she will induce an upward revision of the market estimate of her own ability θ^i . By influencing z_1^j , worker i can bias the learning process in her own favor, regardless the fact that indirectly the estimate of θ^j is also affected. Lemma 1 will help us compare career concerns of permanent and temporary workers.

Lemma 1 *If $\frac{(1+h_j)k_i^2}{(\Sigma_2^{ij} - h_j \Sigma_2^{jj})y_2^{ii'} + (h_j \Sigma_2^{ij} - \Sigma_2^{ii})k_i^2 g_2^{ii'}} > r_i$, then $h_j \beta_2^{i*} + \gamma_2^{i*} > 0$.*

Proposition 4 compares the intensity of reputational incentives of a worker i when she is paired with her colleague at a temporary or permanent basis. What plays a key role is whether her colleague has reputational incentives to help her (when $\rho_1^{ji} > 0$), or sabotage her (when $\rho_1^{ji} < 0$).

Proposition 4 (Career concerns of temporary workers) *Suppose that worker i is low risk-averse so that $h_j \beta_2^{i*} + \gamma_2^{i*} > 0$ holds from Lemma 1.*

(i) *A temporary worker i has stronger reputational incentives to work than a permanent worker i , $M_T^{ii} - M_P^{ii} > 0$, if and only if $\rho_1^{ji} > 0$.*

(ii) *Reputational incentives to help or sabotage arise for a temporary worker i and are always stronger than for a permanent worker i , $|M_T^{ij}| - |M_P^{ij}| > 0$.*

A permanent low-risk averse worker, so that $h_j \beta_2^{i*} + \gamma_2^{i*} > 0$, who receives help from her colleague has weaker reputational incentives to work compared to a temporary worker. Her permanent colleague's help will allow her to manipulate the learning process more effectively. However, the opposite holds when worker i 's higher work effort, e_1^i , induces a downward revision of the estimate of θ^j , implying that worker j will have reputational incentives to sabotage. Worker i now finds it harder to perform well in her own project. Thus, to secure a higher reputation bonus, a long-term worker i who receives sabotage and is

also going to be paired with the same worker j in the next period, she has to focus and work harder on her own project compared to what a temporary worker does. A permanent low-risk averse worker also has weaker reputational incentives to help and stronger incentives to sabotage her colleague compared to a temporary worker. Influencing market perceptions about a colleague's reputation is important for a long-term worker. In that sense, a long-term worker seems to be less 'altruistic', although she will be paired with her co-worker as long as employment lasts. A permanent high risk-averse worker seems to be more altruistic when she tries to build up her reputation, with stronger incentives to help and weaker incentives to sabotage compared to temporary workers.

The contractual incentives will also be different. When pairing is short, a temporary worker does not bear any intertemporal risk associated with her current colleague's human capital. The first-period equilibrium incentive parameters are

$$\beta_{1T}^i = \frac{1}{\Omega_{1T}^i} - M_T^{ii} - r_i \sigma_i^2 g_{1T}''' (\beta_2^{i*} + h_i \gamma_2^{i*}) \frac{k_i^2 + r_i [\varphi_j^2 + (1 - h_i h_j) \sigma_j^2]}{\zeta_{1T}^i \Omega_{1T}^i}, \quad (12)$$

$$\gamma_{1T}^i = \frac{\Phi_{1T}^i}{\Omega_{1T}^i} - M_T^{ij} - r_i \sigma_i^2 y_{1T}''' (\beta_2^{i*} + h_i \gamma_2^{i*}) \frac{h_i + r_i [h_i \varphi_i^2 - h_j (1 - h_i h_j) \sigma_j^2] g_{1T}'''}{\zeta_{1T}^i \Omega_{1T}^i}, \quad (13)$$

The terms Φ_{1T}^i , Ω_{1T}^i and ζ_{1T}^i are the same to those for the permanent workers replacing y_1''' with y_{1T}''' , and g_1''' with g_{1T}''' .

This analysis suggests that if the workers are high risk-averse and there is high uncertainty about their abilities - for example, they can be junior workers hired at the entry level - so that they have incentives to sabotage, the manager will prefer to sign two-year contracts with them. The manager aims to use explicit contracts as a tool to weaken workers' incentives to sabotage in their attempt to build up their reputation. Sabotage between pairs who stay together is not as intense as for temporary workers, given also that the manager is equally happy by hiring a high reputation worker at a high wage or a low reputation worker at a low wage.

5 Conclusion

This paper analyzes compensation contracts and reputational incentives of co-workers when a worker's individual performance depends on the skills of her colleague. We assume that the support a worker receives depends on both her colleague's effort and innate ability, as it is likely to happen in work groups when peer interactions occur. We show that due to the incentives-insurance trade-off, the manager who fully commits to a life-time income path may provide long-term contracts that induce a high risk-averse agent even to sabotage her colleague. Under full commitment, career concerns do not arise since long-term contracts are offered and signed at the beginning of workers' employment. Thus,

sabotage is induced solely by the explicit contracts, aiming at reducing the insurance that needs to be provided and incentivizing the workers to focus on their own projects. We also argue that the duration of employment is a key determinant of workers' contracts. While for a short-term employment, the manager can motivate a worker to help her colleague, as employment is extended to many periods, the same worker may be induced to sabotage.

This paper also examines short-term contracts which are renegotiated in each period. Career concerns now arise and shape explicit contracts. A worker now has reputational incentives to sabotage her colleague because she wants to manipulate market assessments about her own ability. The manager can now set a negative explicit incentive in her attempt to decrease a worker's desire to sabotage. We show that reputational incentives also arise for temporary workers who will be paired with another worker in the next period. This happens because a worker can shape market assessment about her own ability by influencing her current colleague's production.

The present model can be used as a reference point for future works and extensions. One can consider production functions in which a worker's help effort is multiplied with her colleague's ability. The substitutive relationship between implicit and explicit incentives formulated by Gibbons & Murphy (1992) may be challenged. Dewatripont et al. (1999) examine this relationship in a single-agent model in which a worker's effort is multiplied with her *own* ability. This paper can also be used as a reference point to develop a model in which a worker contributes to multiple projects and her commitment to be involved in each project varies. Given that she is paired with workers of different abilities in each project, one can examine if she has incentives to work in projects where her colleagues are of higher or lower productivity, or in projects of longer duration.

The size of the work group and the degree of heterogeneity between the colleagues' skills are other key determinants of compensation contracts in the presence of career concerns. For instance, in software and microelectronics-based industry, the research groups are small and the duration of a research project is short, while in pharmaceutical and biotechnology, the research groups are large, lacking the ability to break them up into small independent modules. The duration of a research project is also long, since it involves experimentation. One can also enrich this framework by considering different allocations of the bargaining power or allowing for side payments between the workers. Other forms of reputation structures, such as mimicking irrational types (Faingold & Sannikov (2011)), and signaling (Mailath & Samuelson (2001)) which rely on asymmetric information and require the workers to signal their type as a team player, will provide interesting insights in team formation.

References

- Arya, A. & Mittendorf, B. (2011), ‘The benefits of aggregate performance metrics in the presence of career concerns’, *Management Science* **57**(8), 1424–1437.
- Auriol, E., Friebel, G. & Pechlivanos, L. (2002), ‘Career concerns in teams’, *Journal of Labor Economics* **20**(2), 289–307.
- Bar-Isaac, H. (2007), ‘Something to prove: reputation in teams’, *RAND Journal of Economics* **38**(2), 495–511.
- Bar-Isaac, H. & Deb, J. (2014), ‘(good and bad) reputation for a servant of two masters’, *American Economic Journal: Microeconomics* **6**(4), 293–325.
- Bar-Isaac, H. & Hörner, J. (2014), ‘Specialized careers’, *Journal of Economics & Management Strategy* **23**(3), 601–627.
- Bernhardt, D. (1995), ‘Strategic promotion and compensation’, *Review of Economic Studies* **62**(2), 315–339.
- Bilanakos, C. (2013), ‘Career concerns and firm-sponsored general training’, *Research in Economics* **67**(2), 117–132.
- Bonatti, A. & Hörner, J. (2017), ‘Career concerns with exponential learning’, *Theoretical Economics* **12**(1), 425–475.
- Chalioti, E. (2016), ‘Team production, endogenous learning about abilities and career concerns’, *European Economic Review* **85**, 229–244.
- Cisternas, G. (2017), ‘Career concerns and the nature of skills’, *American Economic Journal: Microeconomics*, *Forthcoming*.
- Dewatripont, M., Jewitt, I. & Tirole, J. (1999), ‘The economics of career concerns, Part II: Application to missions and accountability of government agencies’, *Review of Economic Studies* **66**(1), 199–217.
- Dewatripont, M., Jewitt, I. & Tirole, J. (2000), ‘Multitask agency problems: Focus and task clustering’, *European Economic Review* **44**(4), 869–877.
- Effinger, M. R. & Polborn, M. K. (2001), ‘Herding and anti-herding: A model of reputational differentiation’, *European Economic Review* **45**(3), 385–403.

- Faingold, E. & Sannikov, Y. (2011), ‘Reputation in continuous-time games’, *Econometrica* **79**(3), 773–876.
- Fama, E. F. (1980), ‘Agency problems and the theory of the firm’, *Journal of Political Economy* **88**(2), 288–307.
- Ferrer, R. (2010), ‘The effect of lawyers’ career concerns in litigation’, *working paper*, *Universitat Pompeu Fabra*.
- Gibbons, R. & Murphy, K. J. (1992), ‘Optimal incentive contracts in the presence of career concerns: Theory and evidence’, *Journal of Political Economy* **100**(3), 468–505.
- Hahn, V. (2017), ‘Committee design with endogenous participation’, *Games and Economic Behavior* **102**, 388–408.
- Holmström, B. (1982a), *Managerial incentives - A dynamic perspective*, Essays in Economics and Management in Honor of Lars Wahlbeck, Helsinki, Finland: Swedish School of Economics.
- Holmström, B. (1982b), ‘Moral hazard in teams’, *Bell Journal of Economics* **13**, 324–340.
- Holmström, B. (1999), ‘Managerial incentive problems: A dynamic perspective’, *Review of Economic Studies* **66**, 169–182.
- Holmström, B. & Milgrom, P. (1987), ‘Aggregation and linearity in the provision of intertemporal incentives’, *Econometrica* **55**(2), 303–328.
- Holmström, B. & Milgrom, P. (1991), ‘Multitask principal agent analyses: Incentive contracts, asset ownership and job design’, *Journal of Law, Economics, and Organization* **7**, 24–52.
- Itoh, H. (1991), ‘Incentives to help in multi-agent situations’, *Econometrica* **59**, 611–636.
- Itoh, H. (1992), ‘Cooperation in hierarchical organizations: An incentive perspective’, *Journal of Law, Economics, and Organization* **8**, 321–345.
- Jeon, S. (1996), ‘Moral hazard and reputational concerns in teams: Implications for organizational choice’, *International Journal of Industrial Organization* **14**, 297–315.
- Laffont, J. J. & Tirole, J. (1988), ‘The dynamics of incentive contracts’, *Econometrica* **56**(5), 1153–1175.
- Lazear, E. (1989), ‘Pay equality and industrial politics’, *Journal of Political Economy* **97**, 561–580.
- Lazear, E. (1999), ‘Personnel economics: Past lessons and future directions’, *Journal of Labor Economics* **17**(2), 199–236.

- Mailath, G. J. & Samuelson, L. (2001), ‘Who wants a good reputation?’, *The Review of Economic Studies* **68**(2), 415–441.
- Meyer, M. A., Olsen, T. E. & Torsvik, G. (1996), ‘Limited intertemporal commitment and job design’, *Journal of Economic Behavior & Organization* **31**(3), 401–417.
- Meyer, M. A. & Vickers, J. (1997), ‘Performance comparisons and dynamic incentives’, *Journal of Political Economy* **105**(3), 547–581.
- Milgrom, P. & Oster, S. (1987), ‘Job discrimination, market forces, and the invisibility hypothesis’, *Quarterly Journal of Economics* pp. 453–476.
- Mookherjee, D. (1984), ‘Optimal incentive schemes with many agents’, *Review of Economic Studies* **51**, 433–446.
- Mookherjee, D. & Ray, D. (2002), ‘Contractual structure and wealth accumulation’, *American Economic Review* **92**(4), 818–849.
- Mukherjee, A. (2008), ‘Sustaining implicit contracts when agents have career concerns: the role of information disclosure’, *RAND Journal of Economics* **39**(2), 469–490.
- Ortega, J. (2003), ‘Power in the firm and managerial career concerns’, *Journal of Economics and Management Strategy* **12**(1), 1–29.
- Pitchford, R. (1998), ‘Moral hazard and limited liability: The real effects of contract bargaining’, *Economic Letters* **61**(2), 251–259.

A APPENDIX

A.1 Long-term explicit incentives

In the beginning of period 1, given (4), the manager maximizes

$$\begin{aligned}
 L = & \sum_{t=1}^2 \sum_{i=1}^2 E \{ z_t^i - \omega_t^i - \beta_t^i z_t^i - \gamma_t^i z_t^j \} + \sum_{i=1}^2 \lambda_i [\beta_t^i - g'(e_t^i)] + \sum_{i=1}^2 \mu_i [k_i \gamma_t^i - y'(a_t^i)] \\
 & + \sum_{t=1}^2 \sum_{i=1}^2 \xi_i [E \{ \omega_t^i + \beta_t^i z_t^i + \gamma_t^i z_t^j \} - g(e_t^i) - y(a_t^i)] - \frac{r_i}{2} \text{Var}(w_1^i + w_2^i) - \Theta_1^i.
 \end{aligned}$$

Omitting details, the Kuhn-Tucker condition with respect to ω_t^i gives $-1 + \xi_i = 0 \Leftrightarrow \xi_i = 1$, implying that the IR^i constraint binds at the optimum. Then, the Kuhn-Tucker conditions become

$$\begin{aligned} \frac{\partial L}{\partial \lambda_t^i} &= \beta_t^i - g'(e_t^i) \geq 0 \text{ or } \lambda_t^i \geq 0, \lambda_t^i [\beta_t^i - g'(e_t^i)] = 0, \forall i, t \\ \frac{\partial L}{\partial \mu_t^i} &= k_i \gamma_t^i - y'(a_t^i) \geq 0 \text{ or } \mu_t^i \geq 0, \mu_t^i [k_i \gamma_t^i - y'(a_t^i)] = 0, \forall i, t \\ \frac{\partial L}{\partial \beta_t^i} &= -r_i [(\beta_1^i + \beta_2^i) (\sigma_i^2 + h_j^2 \sigma_j^2) + (\gamma_1^i + \gamma_2^i) \Sigma^{ij} + \beta_t^i \varphi_i^2] + \lambda_t^i \leq 0 \text{ or } \beta_t^i \geq 0, \frac{\partial L}{\partial \beta_t^i} \beta_t^i = 0, \forall i, t \\ \frac{\partial L}{\partial \gamma_t^i} &= -r_i [(\beta_1^i + \beta_2^i) \Sigma^{ij} + (\gamma_1^i + \gamma_2^i) (\sigma_j^2 + h_i^2 \sigma_i^2) + \gamma_t^i \varphi_j^2] + k_i \mu_t^i \leq 0 \text{ or } \gamma_t^i \geq 0, \frac{\partial L}{\partial \gamma_t^i} \gamma_t^i = 0, \forall i, t \\ \frac{\partial L}{\partial e_t^i} &= 1 - \lambda_t^i g_t^{i''} - g'(e_t^i) \leq 0 \text{ or } e_t^i \geq 0, \frac{\partial L}{\partial e_t^i} e_t^i = 0, \forall i, t \\ \frac{\partial L}{\partial a_t^i} &= k_i - \mu_t^i y_t^{i''} - y'(a_t^i) \leq 0 \text{ or } a_t^i \geq 0, \frac{\partial L}{\partial a_t^i} a_t^i = 0, \forall i, t \end{aligned}$$

We have $\lambda_t^i = \frac{1-g'(e_t^i)}{g_t^{i''}}$ and $\mu_t^i = \frac{k_i - y'(a_t^i)}{y_t^{i''}}$. Given that $\lambda_t^i > 0$ and $\mu_t^i > 0$, equations (4) hold. Thus, we solve the conditions with respect to β_t^i and γ_t^i , and obtain equations (5). Note that substituting $\gamma_f^i = \Phi_f^i \beta_f^i$ into the condition with respect to β_t^i gives $\lambda_f^i = r_i (\Sigma_f^{ii} + \Omega_f^i \Sigma_f^{ij}) \beta_f^i > 0$. A long-term contract provides the same explicit incentives β_f^i and γ_f^i in each period. A positive γ_f^i requires a positive Φ_f^i , whose denominator (see equation (5)) is positive for all h_i, h_j, k_i and k_j . Thus, its sign depends on the sign of the numerator. It is positive when the condition in Proposition 1 holds.

A.2 Short-term explicit incentives

To find the optimal incentives in period 1, we first need to derive the form of

$$Var \{ \hat{w}_1^i + w_2^i \} = Var \{ \hat{w}_1^i \} + Var \{ w_2^i \} + 2Cov \{ \hat{w}_1^i, w_2^i \}, \quad (14)$$

where

$$Var \{ \hat{w}_1^i \} = Var \{ w_1^i \} + Var \{ \omega_2^{i*} \} + 2Cov \{ w_1^i, \omega_2^{i*} \}. \quad (15)$$

The variance of the first-period wage is given by

$$Var \{ w_1^i \} = (\beta_1^i + h_i \gamma_1^i)^2 \sigma_i^2 + (h_j \beta_1^i + \gamma_1^i)^2 \sigma_j^2 + (\beta_1^i)^2 \varphi_i^2 + (\gamma_1^i)^2 \varphi_j^2. \quad (16)$$

Note that

$$\begin{aligned} & \Theta_2^i - E \{ \beta_2^i z_2^i + \gamma_2^i z_2^j \mid z_1^i, z_1^j \} = \\ &= E \{ (1 + h_i - \beta_2^i - h_i \gamma_2^i) \theta^i - (h_j \beta_2^i + \gamma_2^i) \theta^j - \beta_2^i e_2^i - k_i \gamma_2^i a_2^i - k_j \beta_2^i a_2^j - \gamma_2^i e_2^j \mid z_1^i, z_1^j \} \\ &= M_P^{ii} (z_1^i - \hat{e}_1^i - k_j \hat{a}_1^j) + M_P^{ij} (z_1^j - \hat{e}_1^j - k_i \hat{a}_1^i) - \beta_2^i e_2^i - k_i \gamma_2^i \hat{a}_2^i - k_j \beta_2^i \hat{a}_2^j - \gamma_2^i \hat{e}_2^j, \end{aligned}$$

where M_P^{ii} and M_P^{ij} are given by (6) and (7). Thus, the variance of $\omega_2^{i*}(\beta_2^{i*}, \gamma_2^{i*})$ is

$$Var\{\omega_2^{i*}\} = (M_P^{ii} + h_i M_P^{ij})^2 \sigma_i^2 + (h_j M_P^{ii} + M_P^{ij})^2 \sigma_j^2 + (M_P^{ii})^2 \varphi_i^2 + (M_P^{ij})^2 \varphi_j^2. \quad (17)$$

By (16) and (17), we also have

$$\begin{aligned} Cov\{w_1^i, \omega_2^{i*}\} &= (\beta_1^i + h_i \gamma_1^i) (M_P^{ii} + h_i M_P^{ij}) \sigma_i^2 + (h_j \beta_1^i + \gamma_1^i) (h_j M_P^{ii} + M_P^{ij}) \sigma_j^2 \\ &\quad + \beta_1^i M_P^{ii} \varphi_i^2 + \gamma_1^i M_P^{ij} \varphi_j^2. \end{aligned} \quad (18)$$

To obtain the variance of $\hat{w}_1^i$, let $B_1^i \equiv \beta_1^i + M_P^{ii}$ and $\Gamma_1^i \equiv \gamma_1^i + M_P^{ij}$. Then, by (16), (17) and (18), the variance in (15) becomes

$$Var\{\hat{w}_1^i\} = (B_1^i + h_i \Gamma_1^i)^2 \sigma_i^2 + (h_j B_1^i + \Gamma_1^i)^2 \sigma_j^2 + (B_1^i)^2 \phi_i^2 + (\Gamma_1^i)^2 \phi_j^2,$$

and the covariance of $\hat{w}_1^i$ and w_2^i can be written as

$$Cov\{\hat{w}_1^i, w_2^i\} = (\beta_2^i + h_i \gamma_2^i) (B_1^i + h_i \Gamma_1^i) \sigma_i^2 + (h_j \beta_2^i + \gamma_2^i) (h_j B_1^i + \Gamma_1^i) \sigma_j^2.$$

Therefore, equation (14) takes the form

$$\begin{aligned} Var\{\hat{w}_1^i + w_2^i\} &= [B_1^i + \beta_2^i + h_i (\Gamma_1^i + \gamma_2^i)]^2 \sigma_i^2 + [h_j (B_1^i + \beta_2^i) + \Gamma_1^i + \gamma_2^i]^2 \sigma_j^2 \\ &\quad + (B_1^i)^2 \varphi_i^2 + (\Gamma_1^i)^2 \varphi_j^2 + (\beta_2^i)^2 \varphi_i^2 + (\gamma_2^i)^2 \varphi_j^2. \end{aligned}$$

Provided that the IR_1^i constraint binds and Θ_1^i is zero, the first-period base payment is

$$\omega_1^i = -E\{\beta_1^i z_1^i + \gamma_1^i z_1^j\} + g(e_1^{i*}) + y(a_1^{i*}) - E\{w_2^{i*}\} + g(e_2^{i*}) + y(a_2^{i*}) + \frac{r_i}{2} Var\{\hat{w}_1^i + w_2^i\}. \quad (19)$$

We take the Kuhn-Tucker conditions as in appendix (A.1) and solve the equations:

$$\begin{aligned} [B_1^i + \beta_2^i + h_i (\Gamma_1^i + \gamma_2^i)] \sigma_i^2 + [h_j (B_1^i + \beta_2^i) + \Gamma_1^i + \gamma_2^i] h_j \sigma_j^2 + B_1^i \varphi_i^2 &= \frac{\lambda_i}{r_i} \\ [B_1^i + \beta_2^i + h_i (\Gamma_1^i + \gamma_2^i)] h_i \sigma_i^2 + [h_j (B_1^i + \beta_2^i) + \Gamma_1^i + \gamma_2^i] \sigma_j^2 + \Gamma_1^i \varphi_j^2 &= \frac{k_i \mu_i}{r_i} \end{aligned}$$

$$1 - B_1^i - \lambda_i g_1^{i'''} = 0$$

$$k_i (1 - \Gamma_1^i) - \mu_i g_1^{i'''} = 0$$

The optimal β_1^{i*} and γ_1^{i*} are given by equations (8) and (9).

A.3 Temporary workers

To derive the optimal incentives of temporary workers in period 1, we first need to derive the variance $Var\{\hat{w}_1^i + w_2^i\}$ given by (14), while (15) and (16) hold. Suppose that worker i will be paired

with worker l in period 2, whose ability θ^l is normally distributed with zero mean and variance σ_l^2 . Note that

$$\begin{aligned} \Theta_2^i - E \{ \beta_2^i z_2^i + \gamma_2^i z_2^l \mid z_1^i, z_1^j \} &= \\ &= E \{ (1 + h_i - \beta_2^i - h_i \gamma_2^i) \theta^i - (h_l \beta_2^i + \gamma_2^i) \theta^l - \beta_2^i e_2^i - k_i \gamma_2^i a_2^i - k_l \beta_2^i a_2^l - \gamma_2^i e_2^l \mid z_1^i, z_1^j \} \\ &= M_T^{ii} (z_1^i - \hat{e}_1^i - k_j \hat{a}_1^j) + M_T^{ij} (z_1^j - \hat{e}_1^j - k_i \hat{a}_1^i) - \beta_2^i \hat{e}_2^i - k_i \gamma_2^i \hat{a}_2^i - k_l \beta_2^i \hat{a}_2^l - \gamma_2^i \hat{e}_2^l, \end{aligned}$$

where M_T^{ii} and M_T^{ij} are given by equations (10) and (11). Thus, the variance of $\omega_2^i (\beta_2^{i*}, \gamma_2^{i*})$ is

$$Var \{ \omega_2^{i*} \} = (M_T^{ii} + h_i M_T^{ij})^2 \sigma_i^2 + (h_j M_T^{ii} + M_T^{ij})^2 \sigma_j^2 + (M_T^{ii})^2 \varphi_i^2 + (M_T^{ij})^2 \varphi_j^2. \quad (20)$$

By (16) and (20), we also have

$$\begin{aligned} Cov \{ \omega_1^i, \omega_2^{i*} \} &= (\beta_1^i + h_i \gamma_1^i) (M_T^{ii} + h_i M_T^{ij}) \sigma_i^2 + (h_j \beta_1^i + \gamma_1^i) (h_j M_T^{ii} + M_T^{ij}) \sigma_j^2 \\ &\quad + \beta_1^i M_T^{ii} \varphi_i^2 + \gamma_1^i M_T^{ij} \varphi_j^2. \end{aligned} \quad (21)$$

Let $b_1^i \equiv \beta_1^i + M_T^{ii}$ and $g_1^i \equiv \gamma_1^i + M_T^{ij}$. Then, by (16), (20) and (21), the variance in (15) becomes

$$Var \{ \hat{w}_1^i \} = (b_1^i + h_i g_1^i)^2 \sigma_i^2 + (h_j b_1^i + g_1^i)^2 \sigma_j^2 + (b_1^i)^2 \varphi_i^2 + (g_1^i)^2 \varphi_j^2,$$

and the covariance of $\hat{w}_1^i$ and w_2^i is

$$Cov \{ \hat{w}_1^i, w_2^i \} = (\beta_2^i + h_i \gamma_2^i) (b_1^i + h_i g_1^i) \sigma_i^2.$$

Thus, equation (14) takes the form

$$\begin{aligned} Var \{ \hat{w}_1^i + w_2^i \} &= [b_1^i + \beta_2^i + h_i (g_1^i + \gamma_2^i)]^2 \sigma_i^2 + (h_l \beta_2^i + \gamma_2^i)^2 \sigma_l^2 + (h_j \beta_1^i + \gamma_1^i)^2 \sigma_j^2 \\ &\quad + (b_1^i)^2 \varphi_i^2 + (g_1^i)^2 \varphi_j^2 + (\beta_2^i)^2 \varphi_i^2 + (\gamma_2^i)^2 \varphi_j^2. \end{aligned} \quad (22)$$

Given also that the first-period base payment is as in equation (19) and the IR_i binds at the optimum, we take the Kuhn-Tucker conditions as in appendix (A.2) and solve the equations:

$$\begin{aligned} [b_1^i + \beta_2^i + h_i (g_1^i + \gamma_2^i)] \sigma_i^2 + (h_j \beta_1^i + \gamma_1^i) h_j \sigma_j^2 + b_1^i \varphi_i^2 &= \frac{\lambda_i}{r_i} \\ [b_1^i + \beta_2^i + h_i (g_1^i + \gamma_2^i)] h_i \sigma_i^2 + (h_j \beta_1^i + \gamma_1^i) \sigma_j^2 + g_1^i \varphi_j^2 &= \frac{k_i \mu_i}{r_i} \end{aligned}$$

$$\begin{aligned} 1 - b_1^i - \lambda_i^i g_1^{i'''} &= 0 \\ k_i (1 - g_1^i) - \mu_i^i y_1^{i'''} &= 0 \end{aligned}$$

Equations (12) and (13) give the optimal β_{1T}^{i*} and γ_{1T}^{i*} .