

The Dynamics of Exploiting Overconfident Workers*

Matthias Fahn[†] and Nicolas Klein[‡]

September 27, 2024

Abstract

This paper studies a long-term employment relationship with an overconfident worker who updates his beliefs using Bayes' rule. Once the worker has proven to be a good match, exploitation opportunities disappear. Then, it may be optimal to either end the relationship or promote/transfer the worker to a different role, especially if the new position offers fresh opportunities to exploit his overconfidence. In doing so, we offer a novel microfoundation for the “Peter Principle,” rooted in this dynamic of overconfidence exploitation. Our analysis addresses key limitations in previous explanations, particularly those related to the findings of Benson et al. (2019), where the Peter Principle was observed among highly confident workers.

*A previous version of this paper was circulated under the title “Non-Common Priors, Incentives, and Promotions: The Role of Learning.” We thank Miguel Ballester, Gabriel Carroll, Raphael Godefroy, Tilman Klumpp, Rumen Kostadinov, Aditya Kuvalekar, Daniel Krähmer, Stephan Lauermann, Arijit Mukherjee, Takeshi Murooka, Georg Nöldeke, Michel Poitevin, Sven Rady, Heiner Schumacher, Peter Schwardmann, Masha Titova, and Hannu Vartiainen, and seminar audiences at Stony Brook 2022 and the Society for the Advancement of Economic Theory (SAET), for helpful comments.

[†]JKU Linz, CESifo and IZA, matthias.fahn@jku.at.

[‡]Université de Montréal and CIREQ, kleinnic@yahoo.com.

1 Introduction

Humans systematically overestimate their abilities. Many think they are better drivers than the average, more intelligent, or better at predicting political outcomes (Myers, 2010; Bondt and Thaler, 1995; see Meikle et al., 2016 or Santos-Pinto and de la Rosa, 2020 for excellent overviews). Recent evidence points towards the prevalence of such “overconfidence” also in the workplace – among managers (Malmendier and Tate, 2005; Malmendier and Tate, 2015; Huffman et al., 2022) as well as non-executives (Hoffman and Burks, 2020).

We are just beginning to understand the extent and persistence of workers’ overconfidence, and how it may affect the structure of long-term employment relationships. Whereas some studies argue that it can be cheaper for firms to hire overconfident workers who overestimate their chances of achieving a successful outcome (Santos-Pinto, 2008; de la Rosa, 2011; Sautmann, 2013), their focus is on one-shot interactions. But the relevance of such “exploitation contracts” relies on their ongoing use over an extended period of time. If workers learn and update their assessments (as studies such as Grossman and Owens, 2012 or Yaouanq and Schwardmann, 2022 indicate), their exploitation may quickly become infeasible.

In this paper, we show that a firm’s exploitation of a worker’s overconfidence about his talent can *intensify* over time, even though he incorporates informative signals and updates beliefs using Bayes’ rule. This implies that the firm’s expected profits can rise as bad signals about the worker’s talent accumulate and firm and worker become increasingly pessimistic. We apply these results to a firm’s job assignment and promotion decisions and demonstrate that it can be optimal to base a promotion on success in the current job, even if the task requirements in the current and the new job are entirely unrelated. The reason is that a success reduces the uncertainty about the worker’s ability, and a subsequent promotion re-instates belief divergence and consequently exploitation possibilities. Thereby, we provide a microfoundation for the so-called *Peter Principle*, according to which firms prioritize current performance in promotion decisions, often overlooking those with the

greatest potential for success in future roles (Benson et al., 2019). In contrast to the prevailing alternative theoretical explanations, our approach does not rely on (parts of) the worker’s performance being unverifiable, and is thus able to rationalize recent evidence by Benson et al., 2019 for the existence of the Peter Principle among highly confident sales agents whose performance can easily be verified.

Our results are derived in a continuous-time setting, where a risk-neutral principal can hire a risk-neutral agent to work on a task. The agent’s value to the principal contains a deterministic part and his stochastic talent (or match quality), which is either high or low. If talent is high, the agent generates a verifiable extra profit to the principal with some probability at each instant in time. If talent is low, the extra profit is never generated. The agent’s talent is initially uncertain, and both players adjust their beliefs using Bayes’ rule:¹ Once the extra profit materializes for the first time, beliefs of the agent being talented jump to 1. Otherwise, beliefs go down. The agent is *overconfident* about his talent, i.e., his starting belief of being talented exceeds the principal’s.

We derive the following results. First, as long as the agent is (strictly) overconfident, the **optimal compensation contract** offers no payment unless there is a success. From the agent’s perspective, this wage (in expectation) compensates him for his outside option. However, from the perspective of the principal (who holds a smaller belief in the agent’s talent), the expected value of the wage is less than the agent’s outside option. As a result, this setup constitutes an *exploitation contract*. Additionally, the more overconfident the agent is, the less the principal expects to pay. Once the agent reveals high talent through an initial success, it becomes optimal to provide a fixed wage that covers his outside option.

Second, the **expected compensation dynamics** are driven by the evolution of the ratio between the principal’s and the agent’s beliefs. As time passes without success, both beliefs decline and approach zero, causing the wage paid after the first success to increase. However, the principal’s belief

¹Evidence for employer learning is provided by Lange (2007) or Kahn and Lange (2014).

decreases more rapidly than the agent's, leading to a reduction in the agent's expected compensation from the principal's perspective and *increasing the level of exploitation*, even as the agent becomes more pessimistic. Meanwhile, the total profits from employing the agent also contain the extra profit in case he is talented, and this component decreases over time if no success occurs. The balance between the (expected) profits from exploiting the agent and the extra profit if the agent is talented depends on the size of this extra profit and the initial belief gap, reflecting the agent's overconfidence.

This determines our third set of results, the principal's **hiring and firing policy**. There, the principal faces a trade-off between getting a certain payoff (her outside option) when not hiring the agent and a risky payoff associated with experimentation that she obtains in addition to some certain payoff when hiring the agent. In our setting, the value of experimentation stems not only from the possibility that the agent is talented, but also from the gains the principal can make by exploiting the agent's overconfidence. If the expected value of experimentation is relatively low – due to a small (expected) value of the agent's talent and a small difference in beliefs – the principal never hires the agent. Conversely, if the expected value of experimentation is high – driven by a high talent potential or significant exploitation opportunities – the principal will always hire the agent.

When the extra profit is large but the initial belief gap is small, the principal will only hire the agent if she has a sufficiently strong belief that the agent is talented. Here, exploitation opportunities are minimal, and hiring is based primarily on the agent's potential talent. In this case, the principal's value increases with her belief, and the agent is fired after a long enough string of failures. On the other hand, if the extra profit is small but the initial belief gap is large, the principal hires the agent only when she is *sufficiently pessimistic* about his talent. In this scenario, despite expecting the agent to have low talent, the principal benefits from exploiting the large belief difference. Here, the principal's value decreases as her belief increases, meaning her profits can *rise with the agent's failures*, and the agent is fired after a success.

After deriving these results, we discuss some implications. First, since the agent is only paid after a success, our setting seems to predict substantial pay for performance. Indeed, there is evidence that performance-based pay is often observed in occupations where overconfidence is common – such as sales or management. Additionally, we argue that the high bonus could instead manifest itself in a high fixed wage that is continuously paid after an initial success. Second, we discuss the implications of the agent being risk averse instead of being protected by limited liability. We argue that, while the optimal compensation scheme then also contains some fixed wage to limit the agent’s exposure to risk, a bonus conditional on success is still paid to exploit the agent’s overconfidence. Moreover, the total compensation paid after a success is higher for lower beliefs, and the expected costs of hiring the agent (from the principal’s perspective) can decrease in the absence of a success, but only if beliefs are sufficiently high. Finally, we relax our baseline assumption that the principal has full bargaining power. There, we conduct comparative statics on the agent’s outside option – as one means to capture varying degrees of labor-market competition – and show that a higher outside option (which may be caused by a reduced labor supply) increases the chances of being in the case where the agent is fired after being revealed as the high type. Moreover, we consider Bertrand competition for the agent and demonstrate that the hiring decisions in this case are the same as in our main setting where the principal has full bargaining power.

Next, we argue that the principal’s outside option may not only correspond to a termination of the employment relationship, but can also involve her value of assigning the agent to a different position. This reassignment might be a promotion, in particular for the case where the principal’s profits increase with failures and she terminates/reassigns the agent after a success. The reason is that then, the move in positions comes with a high bonus that – as discussed before – can also take the form of a rent the agent accrues after the reassignment. With this interpretation, it might be optimal to promote the agent after a success in the original job, even if the agent’s talents in both jobs are entirely unrelated. In general, the agent’s overconfidence leads the principal to put less weight on the agent’s inherent ability for the new

job than what productive efficiency would call for. This result is further exacerbated if the agent is also overconfident in the second job. Then, a first-job success wipes out the principal’s exploitation opportunities there. Promoting him to the second job again re-introduces uncertainty regarding the agent’s talent, thus creating new room to exploit his overconfidence. Moreover, a worker who is currently not successful but who is expected to be talented in the second stage may instead not be promoted because his continued lack of success increases the firm’s profits from exploiting him in his current job.

This mechanism encourages a policy in which it is not necessarily the best-suited agent who is promoted. Thereby, we provide a micro-foundation for the Peter Principle which, according to Benson et al. (2019), implies that firms prioritize current performance in promotion decisions instead of promoting the ones with the best potential for the new job. In contrast to the alternative explanations we are aware of, our approach can generate the Peter Principle even if the agent’s performance is verifiable. Indeed, Benson et al. (2019) demonstrate that the promotion of sales workers is to a larger extent determined by their verifiable sales than would be justified by their fit for a managerial position. Moreover, this link between sales and promotion is especially strong for so-called “lone wolves” who are highly self-confident but whose fit for managerial positions is particularly poor because of a lack of willingness to collaborate with others.

Related Literature

We contribute to the literature on incentive contracts with overconfident agents. DellaVigna and Malmendier (2004) and Heidhues and Köszegi (2010) provide early work on how to design incentive contracts when consumers are overconfident, in this case about their future self control. They show that exploitation is optimal and feasible. In a static employment setting with a risk-neutral principal and a risk-averse agent, Santos-Pinto (2008) and de la Rosa (2011) demonstrate that implementing effort can be cheaper if

the agent is overconfident about his ability. Moreover, exploitation contracts can emerge, in which an agent’s overconfidence gives him a realized expected utility that is smaller than anticipated by himself. Schumacher and Thysen (2022) explore the consequences of an agent having misspecified beliefs that pertain to the consequences of his actions off the equilibrium path. This can also make it cheaper to provide incentives for a risk-averse agent who underestimates the benefits of shirking.

There also is evidence for the existence of exploitation contracts, in the lab as well as in the field. In the lab, Sautmann (2013) finds that agents who are overconfident about their abilities overestimate their expected payoffs and consequently are worse off than underconfident agents. Larkin et al., 2012 observe that participants who overestimate their performance in a standard multiplication task are more likely to select convex (instead of linear) incentive schemes that offer generous rewards for levels of performance they are unlikely to attain.

Evidence from the field is mostly based on executive compensation, where overconfident managers receive incentive-heavy compensation contracts (Humphery-Jenner et al., 2016). Firms benefit from these arrangements because overconfident CEOs receive fewer bonus payments and smaller stock option grants than their peers and therefore ultimately receive less total compensation (Otto, 2014).

Although the mechanisms underlying such “exploitation contracts” seem well understood, their benefits for firms depend on whether they can repeatedly be applied over a sufficiently long time horizon. Thus, it is important to understand how employees assess the feedback they receive about their performance. If they learn and update their assessments (such as in Yaouanq and Schwardmann, 2022), one might expect their exploitation to quickly become infeasible. We show, however, that learning about the source of the underlying overconfidence can actually exacerbate the agent’s exploitation. Moreover, even if complete learning is achieved, firms may re-instate uncertainty – and consequently overconfidence – by promoting the agent. Existing dynamic models with overconfident agents either rely on environments of mis-

specified learning in which success has several determinants and the agent is overconfident about one of them (Heidhues et al., 2018; Heidhues et al., 2021; Hestermann and Yaouanq, 2021; Murooka and Yamamoto, 2021), or assume that the agent assigns probability 1 to one state of the world and therefore does not update when receiving new information (Englmaier et al., 2020).

We also relate to the theoretical literature on the “Peter Principle,” according to which firms prioritize current performance in promotion decisions over potential ability in the new job. We argue that the previous explanations are insufficient to explain this phenomenon in a setting with sales agents as observed by Benson et al. (2019). For example, one explanation having been proposed is that employees may value the signaling role of promotions. Waldman (1984) and DeVaro and Waldman (2012) set up models in which firms privately observe workers’ abilities for the “new” job. Because a promotion provides a signal about this ability to the market, it has to come with a steep wage increase to fend off counteroffers, and the ability threshold above which someone is promoted is higher than without private information. Yet these theories do not predict that the wrong people are promoted, but instead only the very best. In the setting explored by Benson et al. (2019), a promotion indicates a sales worker’s ability for his current job, rather than managerial talent. As a potential explanation, firms may use promotions instead of monetary bonuses to incentivize workers because the latter are more prone to influence activities by workers (Milgrom and Roberts, 1988), an idea formally modeled by Fairburn and Malcomson (2001). These models rely on an effort dimension that is *not* objectively measurable and can therefore be misreported by supervisors. By the same token, in Lazear (2004), firms do observe but a noisy signal of an agent’s talent. In expectation, a high observation will correspond to a high noise term. Firms anticipate this sub-optimal allocation that is due to mean reversion but, given the information they have access to, they cannot avoid the Peter principle. We therefore conclude that, although these theories are able to rationalize the incentive roles of promotions, they are insufficient to explain the observations made by Benson et al. (2019), which are based on an easily verifiable task and highly confident individuals. Instead, we argue that firms might *intentionally* pro-

mote revealed high performers even though they know this is inefficient – as a consequence of the optimal exploitation of overconfident workers.

2 Model

A principal (“she”) and an agent (“he”) interact in continuous time over an infinite horizon. Both parties discount future payoffs at the rate of $r > 0$. At each instant $t \in \mathbb{R}_+$, the principal can either hire the agent or produce herself. If she produces herself in $[t, t + dt)$, she receives a profit flow of $\bar{\pi}dt \geq 0$. If the agent is hired at instant t , he incurs an (opportunity) cost of $cdt > 0$. This opportunity cost may not only capture the utility of working for a different firm (or not working at all), but also the cost of exerting contractible effort. The agent’s time-invariant talent $\theta \in \{0, 1\}$ determines the principal’s profit flow over those time intervals in which the agent is hired. We use continuous time because it allows us to explicitly characterize value functions. Our results below on how the principal’s cost of hiring the agent evolve over time would also apply in discrete time.

Indeed, if the agent is hired at a flow wage of $w \in \mathbb{R}_+$ over a time interval $[t, t + dt)$, the principal’s profit flow over that period is given by $(1 - w)dt + \eta$ with probability θadt , and $(1 - w)dt$ with the counter-probability, for some $a > 0$ and $\eta > 0$. The parameter a thus governs the speed with which a talented agent (i.e., one with $\theta = 1$) produces a breakthrough success (of value η to the principal), and therefore the speed at which the talented agent reveals his type. The principal initially believes that the agent is talented with probability $p_0^P \in (0, 1)$; the agent initially believes that he is talented with probability $p_0^A \in [p_0^P, 1)$. We thus assume that $p_0^A \geq p_0^P$, i.e., the agent is *over-confident*. Both players update their beliefs according to Bayes’ rule: as soon as an extra profit has been observed, both players’ beliefs jump to 1, and stay there. If no extra profit has arrived by period t , party i ’s belief can be written as $p_t^i = \frac{p_0^i e^{-a \int_0^t h_\tau d\tau}}{p_0^i e^{-a \int_0^t h_\tau d\tau} + 1 - p_0^i}$, where we write $h_\tau = 1$ ($h_\tau = 0$) if the agent is (not) hired at instant τ . Let us write beliefs in the form of the odds ratio; in particular, we write $x_t = p_t^A / (1 - p_t^A) = x_0 e^{-a \int_0^t h_\tau d\tau}$, and

$x_t^P = p_t^P / (1 - p_t^P) = x_0^i e^{-a \int_0^t h_\tau d\tau}$. Thus,

$$\frac{x_t^P}{x_t} = \Psi \in (0, 1]$$

is constant over time; Ψ is an inverse measure of the agent's over-confidence, with $\Psi = 1$ corresponding to the case of common priors. In the following, we shall refer to x_t (Ψx_t) as the agent's (principal's) *belief* at instant t . This formulation simplifies our exposition in two dimensions: first, it allows us to focus on just one variable, x_t , to track the evolution of beliefs (instead of p_t^P and p_t^A); second, the agent's overconfidence can be represented by the constant Ψ . However, from time to time we will still refer to the untransformed beliefs p_t^P and p_t^A when it helps to better convey intuition.

Note that Ψ has an additional interpretation. It equals $\lim_{t \rightarrow \infty} p_t^P / p_t^A$ conditional on no success being observed; thus, although each of the beliefs approaches zero in that case, the limit of the ratio is strictly positive. This interpretation will become important when we discuss the dynamics of the costs of hiring the agent.

Contracts, Information, and Equilibrium The principal does not have any long-term commitment power; i.e., she is restricted to offering spot contracts. We furthermore restrict our attention to *Markov spot contracts*. These specify the agent's instantaneous wage payment as a function of the principal's current profit, which is assumed to be verifiable, and the players' current beliefs. The agent is protected by limited liability; i.e., all wage payments must be non-negative at all times.

The agent's belief is common knowledge. We do not need to specify whether the agent is aware of the principal's belief as long as the agent's belief, and his overconfidence, are not affected by the principal's contract offer. For example, both might agree to disagree. We solve for a perfect Bayesian equilibrium (PBE) that maximizes the principal's profits (given her beliefs).

3 Results

First, we derive the optimal compensation structure. It is possible to offer a spot contract that pays the agent his opportunity cost c , independently of beliefs about θ . The agent would be willing to accept such an offer, which would allow the principal to extract the whole rent generated by the agent's employment. However, with $\Psi < 1$, i.e., with $p_0^A > p_0^P$, it is optimal for the principal to exploit the agent's overconfidence and only to pay him conditionally on his producing the extra profit η . The reason is that the agent's belief of being talented and thus of receiving the payment is higher than the principal's, so that both players gain by engaging in a side-bet on the arrival of the extra profit. The risk-neutral agent is willing to accept any contract that at least covers his opportunity cost in expectation, $c/ap_t^A = (1 + x_t)c/ax_t$. In a profit-maximizing equilibrium it is suboptimal to leave the agent a rent. These considerations lead to the following:

Remark 1 *After a success at time t , the principal will pay the agent a lump-sum wage of $W_t = (1 + x_t)c/ax_t$; wages are 0 in the absence of a success. The cost of hiring the agent from the principal's perspective, which in the following we refer to as principal-expected cost, amounts to*

$$\frac{p_t^P}{p_t^A}c = a\Psi \frac{x_t}{1 + \Psi x_t}W_t = \frac{1 + x_t}{1 + \Psi x_t}\Psi c.$$

Note that the principal-expected cost of hiring the agent is smaller than c^2 and, for a given x_t , is increasing in Ψ . Thus, the greater the agent's overconfidence (and thus the lower Ψ), the lower the amount the principal expects to pay the agent for his services.

This structure is (strictly) optimal (for $\Psi < 1$) as long as there has been no success. Once the extra profit has been realized and both players' beliefs jump to 1, this contract generates the same profits as one in which the principal just pays a flow of c irrespectively of whether η materializes or not.

²The optimality of such side-bets is widely known in settings with non-common priors, see Eliaz and Spiegler (2006), or Grubb (2015) for an overview. See also Santos-Pinto (2008) for a risk-averse agent.

3.1 The Cost of Learning

Now, we explore how the agent's expected compensation evolves over time. Clearly, after a success, beliefs jump to 1 and stay there forever thereafter, which implies that expected hiring costs then also become time-invariant. As long as no success has been realized, though, these expected costs decrease as time passes.

Lemma 1 *W_t is decreasing in x_t and hence increasing in time t if there is no success.*

The principal-expected cost of hiring the agent,

$$\frac{p_t^P}{p_t^A}c = a\Psi \frac{x_t}{1 + \Psi x_t} W_t = \frac{1 + x_t}{1 + \Psi x_t} \Psi c,$$

is increasing in x_t ; it tends to Ψc as $x \rightarrow 0$, and to c as $x \rightarrow \infty$. It is a martingale on the principal's information filtration; in case of a success, it jumps up to c , and is decreasing in time t if there is no success.

Without any success, both the principal's and the agent's beliefs go down and eventually approach zero. Because $p_t^P < p_t^A$, though, Bayes' rule indicates that the *relative* reduction of beliefs,

$$\frac{dp_t^-}{p_t} = -a(1 - p_t)dt,$$

is more pronounced for the principal than for the agent. Indeed, on account of the agent's overconfidence, the principal's posterior goes down faster than the agent's. This allows the principal to keep exploiting the agent by promising to offer him an increasingly higher payment for success, which takes place with an ever smaller probability. Hence, as failures accumulate, the agent continues to accept the contract and is exploited every time as his expected compensation decreases.³

This implies that learning does not necessarily benefit the agent. If no success is observed and negative signals accumulate, the agent's (principal-)expected

³We are indebted to an anonymous referee for suggesting this intuition.

compensation goes down. Therefore, even if agents update their beliefs about the underlying source of their overconfidence using Bayes' rule (for which there is evidence, see Yaouanq and Schwardmann, 2022), their exploitation need not vanish in the long run – to the contrary, it may even exacerbate.

Note that this result does not rely on time being continuous but also holds if time is discrete.

3.2 The Optimal Hiring and Firing Decision

The principal's strategy thus boils down to, at each instant, choosing whether to hire the agent as a function of the previous history. As time moves on and no success has been realized, there are 2 countervailing effects on the principal's profits: a direct negative *productivity effect* because the agent is less likely to be talented, and the indirect positive *exploitation effect* (which reflects the evolution of “betting gains” generated by the agent's overconfidence) because incentivizing the agent becomes cheaper. Besides these two myopic effects, the principal's decision can be influenced by benefits of learning about the agent's talent.

Myopic Payoff First, we abstract from learning benefits and derive the conditions under which the productivity effect dominates the exploitation effect, and vice versa. To do so, we set up the principal's myopic (net) payoff of employing the agent,

$$\mathcal{M}(x) := \left[1 + \frac{\Psi x}{1 + \Psi x} a\eta - \frac{1 + x}{1 + \Psi x} \Psi c - \bar{\pi} \right].$$

The myopic payoff $\mathcal{M}(x)$ contains the value of hiring the agent, 1, plus the (principal-)expected value of the extra profit which is solely a function of her own belief $p^P = \frac{\Psi x}{1 + \Psi x}$. The third term, $\frac{p^P}{p^A} c = \frac{1 + x}{1 + \Psi x} \Psi c$, indicates the principal-expected costs of hiring the agent, and the fourth term the opportunity costs of not producing herself.

Many of our results will be driven by whether the myopic payoff increases or decreases in the belief x , i.e., the sign of

$$\mathcal{M}'(x) = \Psi \frac{a\eta - (1 - \Psi)c}{(1 + \Psi x)^2}.$$

This yields

Lemma 2 *$\mathcal{M}(x)$ is strictly increasing if $a\eta - (1 - \Psi)c > 0$, strictly decreasing if $a\eta - (1 - \Psi)c < 0$, and constant if $a\eta - (1 - \Psi)c = 0$.*

The sign of $\mathcal{M}'(x)$ does not depend on the current belief x but only on fundamentals. If the extra benefit $a\eta$ is relatively large, $\mathcal{M}'(x) > 0$. Then, the positive productivity effect dominates the negative exploitation effect, and a higher x increases (myopic) profits. If, to the contrary, $a\eta$ is relatively small and the agent's overconfidence pronounced, i.e., Ψ is small, $\mathcal{M}'(x) < 0$. Then, the negative exploitation effect dominates, and a *smaller* belief x increases (myopic) profits.

Next, we discuss the conditions under which $\mathcal{M}(x) \geq 0$, i.e., when a myopic principal would hire the agent. For this, we compute profits if only failures have been observed and thus beliefs approach zero, $\lim_{x \rightarrow 0} \mathcal{M}(x) = 1 - \bar{\pi} - c\Psi$, and the myopic payoff if the agent is known to be talented, $\lim_{x \rightarrow \infty} \mathcal{M}(x) = 1 - \bar{\pi} + a\eta - c$. In the following, with a slight abuse of notation we shall write $\mathcal{M}(0)$ for the former and $\mathcal{M}(\infty)$ for the latter.⁴

Now, Lemma 2 implies that $\mathcal{M}(0) < \mathcal{M}(\infty)$ if $\mathcal{M}'(x) > 0$, and $\mathcal{M}(0) > \mathcal{M}(\infty)$ if $\mathcal{M}'(x) < 0$. Therefore, a sufficient condition for $\mathcal{M}(x) \geq 0$ for all x is that $\min\{1 - \bar{\pi} - c\Psi, 1 - \bar{\pi} + a\eta - c\} \geq 0$; in this case, a myopic principal would always hire the agent. By the same token, a sufficient condition for $\mathcal{M}(x) \leq 0$ for all x is that $\max\{1 - \bar{\pi} - c\Psi, 1 - \bar{\pi} + a\eta - c\} \leq 0$; in this case, a myopic principal would never hire the agent. If $1 - \bar{\pi} - c\Psi < 0 < 1 - \bar{\pi} + a\eta - c$,

⁴There is a discontinuity in payoffs at $x = 0$, which stems from the fact that, at $x = 0$, the contract we are looking at (payments contingent on success) ceases to be possible. As our contract continues to be possible, and (weakly) optimal, when $p^A = p^P = 1$, there is no such discontinuity at $x = \infty$.

a myopic principal would hire the agent if and only if $x \geq -\frac{1-\bar{\pi}-\Psi c}{\Psi(1-\bar{\pi}+a\eta-c)} =: x^m$. If, however, $1 - \bar{\pi} + a\eta - c < 0 < 1 - \bar{\pi} - c\Psi$, a myopic principal would hire the agent if and only if $x \leq x^m$. We note that $x^m \in (0, \infty)$ in both these cases.

Learning Benefits On top of the myopic payoff, the principal also takes potential learning benefits of employing the agent into account. Note that $\mathcal{M}(x)$ can be written as $1 - \bar{\pi} + p^P a\eta - (p^P/p^A)c$. As both p^P and p^P/p^A are martingales on the principal's information filtration (see Lemma 1), it follows that $\mathcal{M}(x)$ is also a martingale on the principal's information filtration. Therefore, the principal's expected myopic value when committing to *permanently* employ the agent would be determined by today's belief x , i.e., by $\mathcal{M}(x)$. However, after "bad" outcomes the principal has the option to discontinue employment and thereby cut her losses. Therefore, even if myopic profits are (slightly) negative, employing the agent can be optimal if the principal will thereafter continue employment after some, but fire the agent after other, outcomes. Put differently, if there are beliefs for which myopic payoffs are positive and beliefs for which they are negative, learning can generate benefits. The concept of forgoing current payoffs in exchange for information that is then parlayed into better decisions in the future is what the literature commonly refers to as *experimentation*. It implies that the myopic cutoff x^m does not necessarily dictate the principal's hiring decision. Instead she may be hiring the agent even if her current payoffs would be higher if she produced herself.

Optimal Hiring & Firing To relate these insights to our setting, define $V(x)$ as the total (net) value of employing the agent, evaluated at an instant in time. Thus, it equals the total discounted payoff stream multiplied with the discount rate r (we normalize $V(x)$ to attain comparability with $\mathcal{M}(x)$). Therefore, $V(x)$ equals the myopic payoff $\mathcal{M}(x)$ plus potential benefits of learning. Moreover, the principal profit-maximizing value $V^*(x) = \max\{0, V(x)\}$.

Now, if $\mathcal{M}(0) = 1 - \bar{\pi} - \Psi c$ and $\mathcal{M}(\infty) = 1 - \bar{\pi} + a\eta - c$ both are positive, then $\mathcal{M}(x)$ is positive for all x and the principal would always want to hire the agent. In this case, there are no benefits of learning, and $V(x) = \mathcal{M}(x)$. Learning benefits are also absent if $\mathcal{M}(0)$ and $\mathcal{M}(\infty)$ both are negative. Then, the principal would never want to hire the agent, and $V^*(x) = 0$ for all x .

These (and some additional) results are collected in the following proposition.

Proposition 1 *The subsequent cases describe the conditions for always or never hiring the agent being optimal.*

- [1.] If $\min\{1 - \bar{\pi} - c\Psi, 1 - \bar{\pi} + a\eta - c\} \geq 0$, the principal hires the agent for all $x \in \mathbb{R}_+ \cup \{\infty\}$. The value function is given by $V^*(x) = \mathcal{M}(x) = 1 - \bar{\pi} + \frac{\Psi x}{1 + \Psi x} a\eta - \frac{1 + x}{1 + \Psi x} c\Psi$. If $a\eta > (1 - \Psi)c$, it is strictly increasing and strictly concave; if $a\eta < (1 - \Psi)c$, it is strictly decreasing and strictly convex.
- [2.] If $\max\{1 - \bar{\pi} - c\Psi, 1 - \bar{\pi} + a\eta - c\} \leq 0$, the principal does not hire the agent for any $x \in \mathbb{R}_+ \cup \{\infty\}$. The value function is $V^* = 0$ in this case.

The principal faces a trade-off between getting a sure payoff from not hiring the agent and a risky payoff associated with experimenting by hiring an agent of uncertain talent. When the value of experimentation is low (because the agent's talent is not very important to the principal's production process and there are not many exploitation benefits because the difference in beliefs is modest), the principal prefers never to hire the agent. If, by contrast, the expected value of experimentation is high (because the agent's talent is important to the principal and there are large exploitation gains on account of a large difference in beliefs), the principal always hires the agent.⁵

⁵We are indebted to an anonymous referee for suggesting this intuition.

Now we explore the consequences of $\mathcal{M}(0)$ and $\mathcal{M}(\infty)$ having different signs, in which case there will be learning benefits. For the following, we define $x^* := \frac{r}{r+a}x^m$ and $\tilde{x} := \frac{r+a}{r}x^m$; clearly, $x^* < x^m < \tilde{x}$, where x^m is the myopic cutoff. As before, our main results will depend on whether $\mathcal{M}(0) = 1 - \bar{\pi} - \Psi c$ is smaller or larger than $\mathcal{M}(\infty) = 1 - \bar{\pi} + a\eta - c$ (and thus on the sign of $\mathcal{M}'(x)$). We now show that the principal's hiring decision will admit of a simple cutoff structure, in that she will hire the agent if x is either above or below a certain cutoff, depending on the importance of the extra benefit and the extent of the agent's overconfidence.

First, assume $\mathcal{M}'(x) > 0$ and that $1 - \bar{\pi} + a\eta - c > 0 > 1 - \bar{\pi} - c\Psi$. This implies not only that the positive productivity effect of a higher x dominates the negative exploitation effect, but also that the myopic profit is positive if the belief is sufficiently high ($x \geq x^m$) but negative if the belief is low ($x < x^m$). Then, the principal will hire the agent if and only if she is optimistic enough about his talent, as the following proposition shows.

Proposition 2 *Assume $1 - \bar{\pi} + a\eta - c > 0 > 1 - \bar{\pi} - c\Psi$. Then, the principal hires the agent if and only if $x > x^*$. In this range, V^* is strictly increasing.*

Proofs and closed-form solutions of the value functions are provided in the Appendix.

If $1 - \bar{\pi} + a\eta - c > 0 > 1 - \bar{\pi} - c\Psi$, the principal is mostly interested in the agent's talent, rather than in her exploitation opportunities. In this case, if $x_0 > x^*$, the principal will initially hire the agent and keep hiring him until the belief reaches x^* (if $x_0 \leq x^*$, the principal will never hire the agent). As soon as a success is observed, the agent is hired forever.⁶ x^* is smaller than the myopic cutoff, x^m , because of the benefits of learning. These make it optimal to hire the agent even if the myopic profits are (slightly) negative.

⁶This case is equivalent to a standard one-armed Poisson bandit problem, in which the risky arm is pulled whenever the decision maker is optimistic enough about its quality. The value function in this case is smooth, verifying the usual *smooth pasting* property. As a stylized formalization of the trade-off between experimentation and exploitation, the bandit problem goes back to Thompson (1933) and Robbins (1952). Gittins (1974) showed the structure of the optimal policy; Presman (1991) calculated the *Gittins Index* for the case in which the underlying uncertainty is modeled by a Poisson process.

Second, assume $\mathcal{M}'(x) < 0$ and that $1 - \bar{\pi} + a\eta - c < 0 < 1 - \bar{\pi} - c\Psi$. This implies not only that the negative exploitation effect of a higher x dominates the positive productivity effect, but also that the myopic profit is positive if the belief is sufficiently low ($x \leq x^m$) and negative if the belief is high ($x > x^m$).

Then, the principal will hire the agent if and only if she is *pessimistic* enough about his talent, as the following proposition details.

Proposition 3 *Assume $1 - \bar{\pi} + a\eta - c < 0 < 1 - \bar{\pi} - c\Psi$. Then, the principal hires the agent if and only if $x \leq \tilde{x}$. In this range, V^* is strictly decreasing.*

If $1 - \bar{\pi} + a\eta - c < 0 < 1 - \bar{\pi} - c\Psi$, the principal is less interested in the agent's talent than she is in exploiting him. In this case, if $x_0 \leq \tilde{x}$, the principal will hire the agent until he produces the extra profit, at which time she will permanently fire him (if $x_0 > \tilde{x}$, the principal will never hire the agent). If no success is observed, the agent is hired forever.

Finally, the following remark indicates that no matter if $V^*(x)$ is increasing or decreasing, the principal hires the agent more if the extent of the latter's overconfidence is larger.

Remark 2 *Less similar beliefs (smaller Ψ), and therefore more exploitation opportunities, lead to more experimentation. Thus:*

- In Proposition 2, $\frac{\partial x^*}{\partial \Psi} > 0$.
- In Proposition 3, $\frac{\partial \tilde{x}}{\partial \Psi} < 0$.

We end this section by collecting the monotonicity results for the value function.

Remark 3 *The value function V^* is monotonically increasing if and only if $a\eta \geq (1 - \Psi)c$; it is constant if and only if $a\eta = (1 - \Psi)c$. It is monotonically decreasing if and only if $a\eta \leq (1 - \Psi)c$.*

3.3 Discussion

3.3.1 Performance Pay and Overconfidence

The agent’s overconfidence makes it optimal to pay the agent only after success. Thus, empirically, our mechanism would seem to generate substantial pay for performance. However, most workers work in industries where performance pay is only a small fraction of compensation (Lemieux et al., 2009). Therefore, the question is whether our mechanism applies primarily in the sector of labor markets that can be empirically identified by the presence of substantial pay for performance. We would argue that this is only partially true.

On the one hand, Lemieux et al. (2009) indeed find that sales jobs have the highest incidence of pay for performance, followed by managers⁷. One reason for this is the relative ease of verifying performance in these roles. However, these occupations are also known for widespread overconfidence, which further supports the advantages of pay-for-performance systems. Additionally, substantial evidence suggests that financial market professionals, such as traders and investment bankers, tend to be overconfident in their knowledge of financial markets or their ability to forecast stock prices (Puetz and Ruenzi, 2011; Glaser et al., 2012; Menkhoff et al., 2013), providing additional support for the link between the prevalence of performance pay and overconfidence.

On the other hand, the optimal structure of the compensation scheme – in which the first success generates the highest payment, especially if it took a long time to materialize – can also be interpreted in the following way. First, if $\mathcal{M}'(x) < 0$ and $\mathcal{M}(\infty) < 0$, the negative exploitation effect dominates the positive productivity effect, and the agent is fired after a success and after receiving a substantial payment. This can be interpreted as a severance payment, which would then increase over the agent’s tenure. Second, we argue below that the agent may be reassigned/promoted after a success.

⁷See Malmendier and Tate (2005), Goel and Thakor (2008), Gervais et al. (2011), Malmendier and Tate (2015), for evidence on overconfidence among managers.

Instead of a big bonus upon promotion, the compensation could take the form of a fixed wage that is constantly paid in the new position (which would be strictly optimal if the agent is risk averse, as discussed in the next section). Note that such a structure would require long-term commitment on the part of the principal, which we rule out. However, reputation mechanisms could serve this purpose, which are likely to be easier to implement and enforce if wages are tied to job titles rather than individual employment histories.

Finally, we would argue that an interesting implication of our (and related) work is to show that pay for performance can be optimal even when it is not necessary to incentivize performance. This argument holds even if the agent is risk-averse, as discussed in the next section.

3.3.2 Risk Aversion

Our results are based on the principal and agent having different priors regarding the agent’s talent, making side bets optimal. Side bets take the form of a “performance-based” compensation that is optimal even if the agent does not need to be incentivized (to exert effort). With risk neutral players, there must be a constraint on side bets because otherwise, players would agree on infinite amounts. In our setting, this constraint is the agent’s limited liability, which implies that the size of W_t – his compensation – decreases in the probability that it will have to be paid out. To the contrary, the compensation the principal *expects* to pay decreases over time in the absence of success. The question is whether these features are a consequence of the limited-liability assumption or if they also emerge under alternative settings.

In the following, we consider risk aversion on the side of the agent, a standard friction in agency models (also with an overconfident agent; see Santos-Pinto, 2008, or de la Rosa, 2011). We discuss to what extent the agent’s risk aversion affects the optimality of side bets, which form they take, and how the agent’s compensation evolves over time. If compensation was designed as in our main model, then – with small x – a high W would be paid with a low probability. This would expose the agent to substantial risk which is

expensive for the principal. Therefore, letting the agent's compensation only be success-based will generally not be optimal and some fixed compensation will be paid as well. Moreover, the overall implications of risk aversion will depend on whether the agent has wealth, whether he has access to borrowing/savings devices, as well as the specific form of his utility function.

A complete analysis with a risk-averse agent is beyond the scope of this paper; here, we focus on discussing the case in which the principal maximizes her myopic payoff. This still allows us to generate insights into how W_t , as well as the principal-expected compensation, and the myopic profits evolve. Also recall that, in our main model, when $\mathcal{M}(x)$ is increasing/decreasing, the same holds for the principal's value.

Now, suppose that over a time interval $[t, t + dt)$, the agent receives a flow utility of $u(w)dt$, with $u' > 0$, $u'' < 0$ and $\lim_{w \rightarrow 0} u'(w) = \infty$. If a success is realized, which happens with probability $\theta a dt$, the agent also receives a bonus b and obtains utility $v(w; b)$, with $v(w; b) = u(w + b) - u(w)$.

Moreover, the agent has no alternative source of consumption, and borrowing/saving are not possible (without these restrictions, risk aversion would matter less and our analysis would be closer to our baseline case). His reservation utility over this time interval is $c dt$.

Therefore, the agent's expected utility if working for the principal is

$$u(w)dt + ap^A dt (u(w + b) - u(w)),$$

and the (PC) constraint becomes

$$[(1 - ap^A)u(w) + ap^A u(w + b)] dt \geq c dt.$$

Our objective is to maximize the principal's myopic profits $\mathcal{M}(x) = \max \{0, (1 + p^P a (\eta - b) - w) dt\}$, subject to the (PC) constraint. This yields

Proposition 4 *Assume the agent is risk averse as specified above. Then, if the agent is hired, the profit-maximizing compensation scheme is characterized by the following optimality conditions:*

$$\begin{aligned} [1 + (1 - a) \Psi x] u'(w + b) - \Psi [1 + (1 - a) x] u'(w) &= 0 \\ axu(w + b) + [1 + (1 - a) x] u(w) - (1 + x) c &= 0. \end{aligned}$$

Therefore, $w > 0$ for all x ; $b > 0$ if $x < \infty$ and $\Psi < 1$. If the agent is known to be talented, $b = 0$ is strictly optimal.

Interestingly, if the agent is overconfident, then as long as no success has been realized, using the success-based bonus is always optimal even though the agent is risk averse. Therefore, we argue that overconfidence provides an additional rationale for the use of incentive pay. This is reminiscent of a classic result in portfolio theory, which states that an investor, regardless of their level of risk aversion, should always invest some of his wealth in a risky asset if that asset yields a positive net return.

To get a better idea about potential implications, we now assume $u(W) = \ln W$. Using the optimality conditions derived for Lemma 4, wage and bonus if the agent is hired become

$$\begin{aligned} b &= \left(\frac{1 - \Psi}{\Psi [1 + (1 - a) x]} \right) w \\ \ln(w) &= \frac{(1 + x) c - ax \ln \left(\frac{[1 + (1 - a) \Psi x]}{\Psi [1 + (1 - a) x]} \right)}{1 + x}. \end{aligned}$$

The latter implies that, for $x \rightarrow \infty$, $\ln(w) \rightarrow c$. In the proof to Lemma 4, we also show that $w + b$ decreases in x . Therefore, as with a risk-neutral agent, total compensation goes up over time in the absence of success.

To gain further insights, we assume $a = 0.5$ and $c = 1$, and present some

results for the principal-expected compensation,

$$\begin{aligned} & w + ap^P b \\ &= \frac{(1+x)[1+x\Psi(1-a)]}{(1+\Psi x)[1+(1-a)x]} w. \end{aligned}$$

Then, for $\Psi = 0.2$, the principal-expected compensation has a minimum at $x = 1.17$ (i.e., at $p^A = 0.54$ and $p^P = 0.19$), and is increasing for lower and decreasing for higher values. For $\Psi = 0.5$, the principal-expected compensation has a minimum at $x = 0.82$ ($p^A = 0.45$ and $p^P = 0.29$). For $\Psi = 0.8$, it has a minimum at $x = 0.68$ ($p^A = 0.40$, $p^P = 0.35$).

Therefore, as with risk neutrality, the principal-expected compensation may decrease as long as no success occurs, but only if x is sufficiently high. Then the effect of the reduction in relative beliefs more than compensates for the agent's risk costs.

To assess the evolution of the principal's myopic profit $\mathcal{M}(x)$, we assume that the agent is hired for all x . There, we would have to assume that the base profit from hiring the agent is larger than 1 (or that $\bar{\pi}$ is negative) because, with $u(x) = \ln x$, the principal-expected compensation is always larger than 1 unless Ψ is very small. However, since neither the base profit nor $\bar{\pi}$ interact with x , their size has no effect on the comparative statics conditional on hiring the agent, for which only the term

$$p^P a(\eta - b) - w = \frac{\Psi x}{(1+\Psi x)} a\eta - \frac{(1+x)[1+x\Psi(1-a)]}{(1+\Psi x)[1+(1-a)x]} w$$

is relevant.

We first assume that the payoff of obtaining a success, $\eta = 1$. In this case, for $\Psi = 0.2$, $\mathcal{M}(x)$ increases from $x = 0$ to $x = 2.30$ ($p^A = 0.70$, $p^P = 0.32$), then decreases until $x = 8.68$ ($p^A = 0.90$, $p^P = 0.63$), after which it increases again. For Ψ exceeding ~ 0.23 , $\mathcal{M}(x)$ increases for all x .

If $\eta = 0.1$ and $\Psi = 0.2$, $\mathcal{M}(x)$ increases from $x = 0$ to $x = 1.23$ ($p^A = 0.55$, $p^P = 0.20$), then decreasing until $x = 165.95$ ($p^A = 0.99$, $p^P = 0.97$),

after which it increases again. For $\Psi = 0.5$, $\mathcal{M}(x)$ increases from $x = 0$ to $x = 1.11$ ($p^A = 0.53$, $p^P = 0.36$), then decreases until $x = 21.90$ ($p^A = 0.96$, $p^P = 0.92$), after which it increases again. For Ψ exceeding ~ 0.65 , $\mathcal{M}(x)$ increases for all x .

Therefore, as with risk neutrality (and limited liability), the myopic payoff conditional on hiring the agent may increase in the absence of success, but only for intermediate values of x . For very low x , the cost of the agent's risk aversion is too high, for very high x the belief ratio is too close to 1 to allow substantial gains from side bets.

3.3.3 Competition

In our baseline model, we assume that the principal has full bargaining power and holds the agent to his outside option at all times. In this subsection, we relax this assumption and discuss the resulting implications using the following two approaches. First, we assume that more competition for the agent is reflected in a better outside option for the agent and thus a higher opportunity cost of working for the principal, c . Moreover, if we interpret $\bar{\pi}$ as the principal's outside option, it might also incorporate the difficulties of finding an alternative employee. Then, $\bar{\pi}$ goes down if labor market competition goes up. However, the principal still sets the terms of employment – i.e., we keep maximizing her profits subject to the agent's binding (PC) constraint – which builds on evidence that firms have considerable wage-setting power even in thick labor markets (Manning, 2021, Card, 2022). Second, we allow the agent to make take-it-or-leave-it offers to the principal, for example because several principals engage in Bertrand competition for the agent. We show that our main results continue to hold and additional insights can be generated.

For the first approach, note that comparative statics with respect to $\bar{\pi}$ and c affect the principal's value of hiring the agent. Recall that the principal's expected myopic payoff when hiring the agent – which has a direct positive effect on her value function – equals

$$\mathcal{M}(x) = 1 + \frac{\Psi x}{1 + \Psi x} a\eta - \frac{1 + x}{1 + \Psi x} \Psi c - \bar{\pi}.$$

We have shown that the agent is always hired if $\mathcal{M}(x) > 0$ for all x , and that he is hired for beliefs above or below some threshold if $\mathcal{M}(x) > 0$ for some x (recall that we describe net values, thus the total myopic profit when hiring the agent would be $\mathcal{M}(x) + \bar{\pi}$). It follows that the effect of more competition for the agent (for example caused by a lower labor supply), which increases c and decreases $\bar{\pi}$, is ambiguous on $\mathcal{M}(x)$ and thus on the chances of it being positive. Furthermore, we have shown that a higher x can actually reduce the principal's value of hiring the agent. This holds if the sign of $\mathcal{M}'(x)$ is negative which is equivalent to $c > a\eta / (1 - \Psi)$. Therefore, labor market competition affects the sign of $\mathcal{M}'(x)$ only via c , and $\mathcal{M}'(x)$ is more likely to be negative if c is large. Put differently, a lower labor supply increases our chances of being in the case where profits rise in the absence of a success, and where the agent is fired (or promoted, as discussed in Section 4) after being revealed as the high type.

Next, we consider an equilibrium contract that maximizes the agent's expected utility (according to the agent's assessment) subject to the constraint that the principal achieve an expected profit of at least 0, according to the principal's assessment. Then, as in our main setting, side-bets on the arrival of the extra profit are optimal as long as the agent is overconfident. Therefore, the first payment from the principal to the agent will be made once the first success is realized. Now, this payment needs to satisfy the principal's participation constraint and thus equals

$$\frac{1 - \bar{\pi} + p^P a\eta}{p^P a}.$$

Recall that, in our main model, this payment amounts to $c/p^A a$ to cover the agent's (perceived) opportunity costs of working for the principal.

Now, the agent will expect an amount

$$\frac{p^A}{p^P} (1 - \bar{\pi}) + p^A a\eta.$$

The first term, $p^A (1 - \bar{\pi}) / p^P$ reflects betting gains which – as in our main model where we have used the term exploitation effect to describe its evolution – go up as long as no success has been realized. How the agent's value of being employed by the principal evolves over time then also depends on the size of the productivity effect, i.e., the second term $p^A a\eta$.

To develop an idea about these comparative statics, note that the agent's myopic payoff equals

$$\mathcal{M}_A(x) = \frac{1 + \Psi x}{\Psi(1 + x)}(1 - \bar{\pi}) + \frac{x}{1 + x}a\eta - c,$$

with

$$\mathcal{M}'_A(x) = \frac{a\eta - \frac{(1-\Psi)}{\Psi}(1 - \bar{\pi})}{(1 + x)^2}.$$

Therefore, if $a\eta < (1 - \bar{\pi})(1 - \Psi)/\Psi$, $\mathcal{M}'_A(x) < 0$ and $\mathcal{M}_A(x)$ increases as long as no success has been realized. Note that, as in our main setting, the agent's value function inherits the monotonicity properties of the agent's myopic payoff. Moreover, recall that, if we maximize the principal's profits, the derivative of the principal's myopic profit (and thus of her value function), $\mathcal{M}'(x)$, is negative if $c > a\eta/(1 - \Psi)$. Therefore, in both cases the value is decreasing in x for a small η and/or a small Ψ . The following Proposition 5 presents the results for this subsection and shows that the commonalities between both cases are even more pronounced.

Proposition 5 *Solving for a PBE that maximizes the agent's utility (given his beliefs) yields the following outcomes.*

- If $\min\{1 - \bar{\pi} - c\Psi, 1 - \bar{\pi} + a\eta - c\} \geq 0$, the agent is hired for all $x \in \mathbb{R}_+ \cup \{\infty\}$. The agent's value function is given by $V_A^*(x) = \mathcal{M}_A(x) = \frac{1 + \Psi x}{\Psi(1 + x)}(1 - \bar{\pi}) + \frac{x}{1 + x}a\eta - c$. If $a\eta > \frac{(1 - \Psi)}{\Psi}(1 - \bar{\pi})$, it is strictly increasing and strictly concave; if $a\eta < \frac{(1 - \Psi)}{\Psi}(1 - \bar{\pi})$, it is strictly decreasing and strictly convex.

- If $\max\{1 - \bar{\pi} - c\Psi, 1 - \bar{\pi} + a\eta - c\} \leq 0$, the agent is not hired for any $x \in \mathbb{R}_+ \cup \{\infty\}$. The agent's value function is $V_A^* = 0$ in this case.
- If $1 - \bar{\pi} + a\eta - c > 0 > 1 - \bar{\pi} - c\Psi$, the agent is hired if and only if $x > x^*$, where x^* is the same as in Proposition 2. In this range, V_A^* is strictly increasing.
- If $1 - \bar{\pi} + a\eta - c < 0 < 1 - \bar{\pi} - c\Psi$, the agent is hired if and only if $x \leq \check{x}$, where $\check{x}$ is the same as in Proposition 3. In this range, V_A^* is strictly decreasing.

Proposition 5 states that hiring decisions are the same whether the agent or the principal has full bargaining power, only the ranges for which the respective value functions are increasing/decreasing if the agent is always hired are slightly different. This is because

$$\mathcal{M}_A(\infty) = 1 - \bar{\pi} + a\eta - c = \mathcal{M}(\infty)$$

and, since $\mathcal{M}_A(0) = (1 - \bar{\pi} - \Psi c) / \Psi$,

$$\Psi \mathcal{M}_A(0) = 1 - \bar{\pi} - \Psi c = \mathcal{M}(0),$$

i.e., the thresholds above which the myopic payoffs are positive in the limits $x \rightarrow 0$ and $x \rightarrow \infty$ are identical. Moreover, whether myopic payoffs are increasing or decreasing is independent of x , hence if $\mathcal{M}(\infty) > \mathcal{M}(0)$ and consequently $\mathcal{M}' > 0$, the same holds for $\mathcal{M}_A(x)$.

In conclusion, this section has shown that our results are not caused by the principal having full bargaining power and do not disappear when there is competition for the agent. On the contrary, if we assume that a more competitive labor market increases c and decreases $\bar{\pi}$ (but the principal can still make the employment offer), the range for which the principal's profits decrease in x is expanded.

4 Application – Optimal Job Assignment and the Peter Principle

We have demonstrated that the principal benefits from a divergence between her beliefs and those of the agent. However, once the agent has been successful and is revealed to be competent, their beliefs align, and the principal can no longer benefit from an exploitation contract. If the discrepancy was significant, this may result in the principal opting for her outside option following the agent’s first success. Rather than viewing the outside option as terminating the relationship, we now suggest it could represent reassigning the agent to a different job. In this scenario, we assume that the agent can transition to another position but cannot return to the original one. Due to this "one-way" job rotation, we sometimes refer to such a reassignment as a promotion. This interpretation is further reinforced when the reassignment follows a (first) success, where a large bonus could translate into a higher base salary in the new position (as discussed in Section 3.3.1); then, the reassignment is a move from a lower paying job to a higher paying job, which usually is a feature of a promotion (another typical feature, that multiple agents compete for a promotion, is discussed in Section 4.3 below). We will argue that this interpretation can provide a microfoundation for the well-known “Peter Principle”, according to which workers are promoted to their level of incompetence Peter and Hull (1969) or, more precisely, *firms prioritize current performance in promotion decisions at the expense of promoting the ones with the best potential for the next job* (Benson et al., 2019). Below, we will clearly state how we adapt this definition to our setting.

Now, we will take a closer look at how the agent’s overconfidence can generate a reassignment/promotion policy that is based on success in previous jobs, rather than expected success in the new job; in Section 4.4, we relate it to the evidence provided by Benson et al. (2019). Assume the agent starts out in the first job, which is as described in Section 2. At a time of her choosing, the principal can assign the agent to a second job where his value to the principal

is $\bar{\pi}$; not reassigning him thus entails a flow opportunity cost of $\bar{\pi}$, as before. It is, however, not possible to move the agent back again to the first job. For simplicity, we set the value of firing, or temporarily not employing, the agent in the first job to 0. Importantly, there is no correlation between the jobs regarding the agent's talent for either, and he is (weakly) over-confident concerning the first. Potential overconfidence in the second job is explored in Subsection 4.1.

Clearly, the principal will promote the agent at time $\tau^* = \inf \{t \geq 0 : V^*(x_t) < 0\}$, where $V^*(x_t)$ is the value of employing the agent in the first job net of the value of the outside option $\bar{\pi}$. Generally, our results will depend on whether $a\eta$ is larger or smaller than $(1 - \Psi)c$, i.e., whether $V^*(x)$ is increasing or decreasing (see Remark 3).

As a benchmark, we first define the *efficient* reassignment policy, which maximizes the principal's value whose myopic payoff upon hiring the agent is

$$1 + p^P a\eta - \bar{\pi} - c.$$

The efficient reassignment policy would be selected if the principal and agent were the same person or, as we will assume moving forward, if the agent is not overconfident, i.e., $\Psi = 1$. Under this policy, the likelihood of reassigning the agent increases when no success is observed in the first job. Indeed, if $\Psi = 1$, $a\eta > (1 - \Psi)c$, and V^* is increasing in x . Thus, the agent will either be reassigned after a long enough history of failures in the first job – or right away or never. This is because the longer history of failures makes the opportunity costs of reassigning the agent less severe. As V^* is monotone for common priors, the following Lemma is immediate:

Lemma 3 *Under the efficient reassignment policy, there is a cutoff $\bar{\pi}(x)$ such that the agent is reassigned iff $\bar{\pi} > \bar{\pi}(x)$; moreover, $\bar{\pi}(x)$ is increasing.*

With common priors, the agent is never reassigned after a success because the jobs are uncorrelated, meaning success in the first job does not imply

suitability for the second. In fact, the principal seeks to maximize productive efficiency, balancing the agent’s expected productive value in the second job (which remains constant before a promotion) against the opportunity cost of losing the agent in the first job, which increases with x .

Next, assume that the agent is overconfident, i.e., $\Psi < 1$. Then, the results derived in Propositions 9 and 10 can be used to show

Proposition 6 *There is a cutoff $\bar{\pi}(x, \Psi)$ such that the agent is reassigned iff $\bar{\pi} > \bar{\pi}(x, \Psi)$. $\bar{\pi}(x, \Psi)$ is continuous and strictly increasing in x if and only if $a\eta > (1 - \Psi)c$, strictly decreasing in x if and only if $a\eta < (1 - \Psi)c$, and constant in x if and only if $a\eta = (1 - \Psi)c$.*

For all $x < \infty$, $\bar{\pi}(x, \Psi)$ is continuous and strictly decreasing in the players’ belief alignment Ψ , when the principal’s belief $x \cdot \Psi$ is held constant.

If $a\eta < (1 - \Psi)c$, $V^*(x)$ is decreasing and the agent’s value goes up over time in the absence of a success. Once a success occurs, the principal’s value of keeping the agent in the first job falls because of the eliminated exploitation opportunities. Then, the resulting value reduction increases the relative benefits of a reassignment (i.e., the cutoff $\bar{\pi}(x, \Psi)$ drops) even though the success is *not* informative of the agent’s talent in the second job.

For the reasons outlined above, we will primarily refer to a reassignment following success as a promotion. In this context, a promotion leads to the Peter Principle, which we define as workers being *intentionally and inefficiently* removed from their job in which they have proven to be productive and placed in another for which they have not yet demonstrated their suitability.

This is a variation of the specification used by Benson et al. (2019), where a promotion policy that results in the Peter Principle emphasizes current performance over future potential for the next role. It is important to note that, as long as the agent’s value in the two jobs remains uncorrelated, promoting the agent after a success is always inefficient. In our case, the Peter Principle indeed reflects the firm’s optimal policy when workers are overconfident. In

these instances, the agent is promoted following a success because, once his type is revealed, the value of retaining him in the first job becomes too low for the principal. Below, we will explore additional potential consequences of this policy, such as the possibility of promoting the "wrong" worker.

The question now is under what circumstances the benefits of leveraging overconfidence would outweigh the costs of destroying proven good matches between employees and tasks in real labor markets. According to the condition $a\eta < (1 - \Psi)c$, this occurs when the payoff from the agent's talent, η , is not too high, and Ψ is small, indicating significant overconfidence. In Section 4.4, we argue that sales is a notable example, where successful agents are promoted despite not being the most qualified for managerial roles (significant overconfidence was also found among financial-market professionals; see Section 3.3.1). Furthermore, $a\eta < (1 - \Psi)c$ is more likely when the agent's opportunity costs, for working with the principal, c , are relatively low. As discussed in Section 3.3.3, the size of c could represent labor market competitiveness, with factors like lower labor supply or higher unemployment driving a lower c . Consequently, we would predict that increased competition for workers would amplify the occurrence of the Peter Principle as defined here, although (as far as we are aware) no studies have yet examined this link.

Finally, we discuss the optimal reassignment policy if $a\eta > (1 - \Psi)c$. $V^*(x)$ is increasing and the general pattern is the same as with common priors ($\Psi = 1$). Either the agent is immediately (or never) reassigned, or after many failures in the first job have sufficiently reduced the opportunity costs of a reassignment. Still, the threshold $\bar{\pi}(x, \Psi)$ is higher than with $\Psi = 1$ because the exploitation opportunities in the first job decrease in Ψ (holding the principal's belief $x \cdot \Psi$ constant). Therefore, the optimal reassignment policy is also inefficient.

Finally, the last result of Proposition 6 illustrates the fact that the higher the agent's overconfidence the more valuable he is to the principal in the first job.

4.1 Endogenizing $\bar{\pi}$

Now, we endogenize the agent's value in the second job and assume that his overconfidence can extend to it. Assume that the second job also has the features described in Section 2; there is still no correlation between the agent's talent across both jobs. The details can be found below, in Section 6.2 in the Appendix. As before, the agent's value in the second job remains constant as long as he is not reassigned. Therefore, the same effects as in Section 4 obtain, while introducing the agent's overconfidence in the second job allows for additional comparative statics. The reason is that a reassignment/promotion after a success in the first job re-instates uncertainty and overconfidence, and thus again allows the principal to exploit the agent. Therefore, a lower Ψ in the second job (holding the principal's belief there constant) makes it *ceteris paribus* more likely that the agent is promoted after a first-job success.

4.2 Correlated Jobs and Endogenous Overconfidence

We have demonstrated that the agent's overconfidence can affect promotion decisions and give rise to the Peter Principle. We have assumed that the agent's talent across both jobs is not correlated. However, even with a positive correlation between jobs, the agent's overconfidence induces the principal to put less weight on the agent's talent for the second job than what the pursuit of productive efficiency would require. Indeed, while a success in the first job then increases players' beliefs concerning the agent's talent for the second job, this increase is less pronounced than the increase in the belief about his talent for the first job (unless correlation was perfect). Therefore, with common priors such a success should make a reassignment *less* likely. With non-common priors, however, promotion will become more likely whenever the belief divergence is important enough (i.e., whenever Ψ is low enough), as the success also eliminates exploitation opportunities in the first job.

Moreover, a success in the first job could also by itself increase the agent's overconfidence. For example, assume that the agent overestimates the *correlation* between talent across both jobs. This could be the result of an

inherent bias,⁸ or of the principal’s subterfuge. Then, our results would only require the agent to naively believe the principal’s claim that being successful in the first job is indicative of his potential in the second job. In this case, promoting an agent who has proven to be talented in the first job would again create the additional benefit of being able to exploit his overconfidence in the second job. Importantly, this result would not require the agent to be inherently or initially overconfident – instead his overconfidence would endogenously emerge from a wrong belief that talent in one domain transfers to talent in another.

4.3 Two Agents

Finally, we argue that employing overconfident agents may also lead to the principal putting less weight on an agent’s perceived value in the second job when making the decision as to whom among several agents to promote, as compared to the case in which agents are not overconfident. Assume there is some time T at which the principal wants to promote one out of two agents, $i \in \{1, 2\}$. As in Section 4, let the principal’s value of promoting agent i , $\bar{\pi}_i$, be solely given by his (expected) inherent talent in the second job. Without loss, we assume that $\bar{\pi}_1 \geq \bar{\pi}_2$. To isolate the role of an agent’s overconfidence on the principal’s promotion policy and abstract from differences in the opportunity costs of a promotion, we focus on cases in which the principal’s belief $\Psi_i x_i$ is the same for both agents, while only their Ψ_i might differ. As before, the principal’s optimal policy with $\Psi_1 = \Psi_2 = 1$ is based solely on the agent’s perceived value in the second job. Then, we say that the *right* agent is promoted, which in our case is agent 1. The policy of promoting agent 1 is also adopted if both agents produce a success before time T . However, the following proposition shows that there exist parameters such that the “wrong” agent will be promoted.

⁸For example, the widely observed self-attribution bias, in which people attribute their success to their own abilities instead of just being lucky (see Daniel et al., 1998 or Billett and Qian, 2008 for evidence in the context of managers), could be a factor leading to the agent’s attribution of a first-job success to a general skill that also transfers to other realms.

Proposition 7 *Suppose that agent 1 is more overconfident than agent 2 ($\Psi_1 < \Psi_2$), and suppose that the principal wants to promote one of the agents at a time T at which her beliefs satisfy $\Psi_1 x_{1,T} = \Psi_2 x_{2,T}$. There exist parameters satisfying $\bar{\pi}_1 > \bar{\pi}_2$ and $\Psi_1 < \Psi_2$ such that the principal promotes agent 2.*

If $\Psi_1 < \Psi_2 \leq 1$, agent 1's value is higher in the first job due to his greater overconfidence. If the difference between $\bar{\pi}_1$ and $\bar{\pi}_2$ is small compared to the difference between Ψ_2 and Ψ_1 (for example if agent 2 has succeeded but agent 1 has not), the principal might choose to promote agent 2. This decision arises because, despite agent 1 being better-suited for the second job, his higher overconfidence makes him less expensive to incentivize in the first job.

Collecting the insights from this and the previous subsections, and assuming the agent's perceived value in the second job remains constant, we can conclude that an overconfident agent is more likely to be promoted if he has demonstrated talent in the first job. Conversely, he is less likely to be promoted if he has not performed well, which stands in contrast to the benchmark scenario of common beliefs. Therefore, if workers are overconfident, we would expect to see a positive correlation between current performance and promotion, even when the requirements for the two jobs are entirely unrelated.

4.4 Evidence

Using microdata on sales workers, Benson et al. (2019) find evidence for productive mismatches, as promotion policies put too much weight on current performance, as opposed to perceived fit for the new job. Although sales clearly are a *verifiable* performance measure, high sales are not only rewarded with cash compensation, but also increase a salesperson's chances of being promoted to a managerial position. This policy disregards managerial potential and is costly because it reduces managerial quality (measured as value added to subordinate sales) by 30% compared to a counterfactual where

the ones with the highest managerial potential would be promoted. Benson et al. (2019) discuss a number of potential theoretical explanations for these outcomes which, however, we argue cannot fully rationalize their observations, which are based on an easily verifiable task (see the Related Literature Section above). Instead, we argue that it is not the nature of the job that renders the promotion of successful sales agents (instead of those with the best fit) optimal, but their personal characteristics. Indeed, there is evidence that sales agents are particularly prone to being overconfident. Sevy (2016), in a Forbes blog, argues that, because of the availability of clear performance indicators, sales is an environment that attracts people who want to prove their ability. Those who go for sales care about personal advancement and not about helping a team thrive; this is different in sales *management*, where holding back one’s ego and letting others shine is important.

Moreover, whereas Benson et al. (2019) find that collaboration experience is indicative of better *managerial* performance, so-called “lone wolves,” who never collaborate and are known to be highly self confident (Dixon and Adamson, 2011) are significantly more likely to be promoted to a managerial position.

Finally, Bonney et al. (2020) find that salespeople are more overconfident in their assessment of customer opportunities than sales managers, which is striking because sales managers are typically former salespeople who have been promoted into a new role. This result is consistent with our story, if sales managers are promoted because they have proven to be good salespeople and therefore do a better job of evaluating sales opportunities.

5 Conclusion

We have shown that, in a model in which the principal benefits from employing the agent who can create some extra value if he is talented, the monotonicity of the principal’s value function depends on the agent’s overconfidence Ψ . If the agent’s appraisal of his talent is close to that of the principal, i.e., if Ψ is close to 1, the principal’s value function is increasing

in her belief, and the agent is fired after a long enough streak of failures. If, however, the agent is very overconfident, i.e., if Ψ is low, the principal's value function is *decreasing* in her belief, and the agent is fired *after a success*. As in our model *firing* can be interpreted as promotion to a second, unrelated, job, we provide a novel explanation for the well-documented *Peter principle*: As the agent's type becomes common knowledge after a success, a success makes exploitation contracts impossible; thus, if exploiting the agent's overconfidence is an important part of the principal's objective, she will not want to hire the agent in the current job any longer after a success there, preferring to promote him to another job instead, even if this entails sacrificing some productive efficiency.

We have assumed that a success fully reveals the agent's type, i.e., an untalented agent never produces a success. While this makes our model analytically tractable, we expect our main qualitative conclusions to continue to obtain in a setting in which an untalented agent may also at times, albeit less frequently, produce a success.

6 Appendix

Formal Model Description & Closed-Form Solutions

The principal's strategy boils down to, at each instant, choosing whether to hire the agent as a function of the previous history. Formally, the principal's hiring decisions are a process $\{\chi_t\}_{t \in \mathbb{R}_+}$ that is predictable with respect to the available information, where $\chi_t = 1$ if the agent is hired at instant t , and $\chi_t = 0$ otherwise. Clearly, since the principal is restricted to offering stationary Markov contracts, it is without loss to restrict the principal to choosing a hiring strategy that is also Markovian, i.e., a process $\{\chi_t\}_{t \in \mathbb{R}_+}$ such that $\chi_t = \chi(x_t)$ for all $t \in \mathbb{R}_+$, where $\chi : \mathbb{R}_+ \cup \{\infty\} \rightarrow \{0, 1\}$ is a time-invariant function of beliefs.⁹ In summary, the principal chooses a Markov strategy so as to maximize

$$\begin{aligned} \Pi(x) = \mathbb{E} & \left[\int_0^\infty r e^{-rt} \left(1 - \frac{\Psi x_0}{1 + \Psi x_0} \left(1 - e^{-a \int_0^t \chi(x_\tau) d\tau} \right) \right) \chi(x_t) \right. \\ & \times \left. \left(1 - \bar{\pi} - \frac{1 + x_t}{1 + \Psi x_t} \Psi c + \frac{\Psi x_t}{1 + \Psi x_t} a \left(\eta + \max \left\{ 0, \frac{1 - \bar{\pi} + a\eta - c}{r} \right\} \right) \right) dt \middle| x_0 = x \right], \end{aligned} \quad (1)$$

where the expectation is with respect to the belief process $\{x_t\}_{t \in \mathbb{R}_+}$.

Bellman Equation

We now set up the Bellman equation for the problem. It is given by

$$V^*(x) = \max_{\chi \in \{0,1\}} \chi [\mathcal{B}(x, V^*) + \mathcal{M}(x)],$$

where the myopic payoff from hiring the agent, $\mathcal{M}(x)$, has been introduced in the main text, and $\mathcal{B}(x, V^*)$ captures the benefits from learning about the agent's talent. A myopic principal (i.e., one whose discount rate $r \rightarrow \infty$)

⁹Our payoff-maximizing perfect Bayesian equilibrium will thus be a *Markov perfect equilibrium (MPE)* with players' beliefs as a state variable.

would hire the agent at x if and only if $\mathcal{M}(x) \geq 0$. The same policy would be optimal if the principal did not update her belief regarding the agent's talent (e.g., because the agent's talent is continuously drawn anew). Clearly, as we set out in the main text, $\mathcal{M}(x) \geq 0$ for all x if $\min\{1 - \bar{\pi} - c\Psi, 1 - \bar{\pi} + a\eta - c\} \geq 0$; in this case, a myopic principal would always hire the agent.

By the same token, $\mathcal{M}(x) \leq 0$ if $\max\{1 - \bar{\pi} - c\Psi, 1 - \bar{\pi} + a\eta - c\} \leq 0$; in this case, a myopic principal would never hire the agent. If $1 - \bar{\pi} - c\Psi < 0 < 1 - \bar{\pi} + a\eta - c$, $\mathcal{M}(x) \geq 0$, and a myopic principal would thus hire the agent, if and only if $x \geq -\frac{1 - \bar{\pi} - c\Psi}{\Psi(1 - \bar{\pi} + a\eta - c)} =: x^m$. If, however, $1 - \bar{\pi} + a\eta - c < 0 < 1 - \bar{\pi} - c\Psi$, a myopic principal would hire the agent if and only if $x \leq x^m$. We note that $x^m \in (0, \infty)$ in both these cases.

Yet, a principal that is not myopic also takes the learning benefit of employing the agent into account. This learning benefit amounts to $\frac{1}{r}$ times the infinitesimal generator of the process of posterior beliefs applied to the value function V , and can be written as

$$\mathcal{B}(x, V) := \frac{xa}{r} \left[\frac{\Psi}{1 + \Psi x} (\max\{0, 1 - \bar{\pi} + a\eta - c\} - V(x)) - V'(x) \right].$$

We write $V^*(x) = \max\{0, V(x)\}$, where V satisfies the ODE

$$\begin{aligned} & ax(1 + \Psi x)V'(x) + (r + \Psi x(r + a))V(x) \\ &= r[(1 + \Psi x)(1 - \bar{\pi}) - (1 + x)\Psi c + \Psi xa\eta] + \Psi xa \max\{0, 1 - \bar{\pi} + a\eta - c\}, \end{aligned}$$

which is solved by

$$\begin{aligned} V(x) = & 1 - \bar{\pi} + \frac{\Psi x}{1 + \Psi x} a\eta - c\Psi \frac{1 + x}{1 + \Psi x} \\ & - \mathbb{1}_{\{1 - \bar{\pi} + a\eta - c < 0\}} \frac{a}{a + r} \frac{\Psi x}{1 + \Psi x} (1 - \bar{\pi} + a\eta - c) + C \frac{x^{-\frac{r}{a}}}{1 + \Psi x}, \end{aligned}$$

with C denoting a constant of integration. We furthermore note that¹⁰

$$\lim_{x \downarrow 0} V(x) = 1 - \bar{\pi} - \Psi c;$$

$$\lim_{x \rightarrow \infty} V(x) = (1 - \bar{\pi} + a\eta - c) \left(1 - \mathbb{1}_{\{1 - \bar{\pi} + a\eta - c < 0\}} \frac{a}{a + r} \right);$$

in what follows, we shall write $V(0)$ and $V(\infty)$ respectively for these limits.

If $V(0)$ and $V(\infty)$ have the same sign, the principal's hiring decision under (almost) perfect information will be the same, independently of whether that almost perfect information is positive or negative regarding the agent's talent. It is thus no surprise that the principal will make the same hiring decision for all beliefs, and hence the learning benefit $\mathcal{B} = 0$ in this case, as the following proposition, which restates Proposition 1 from the main text, shows.

Proposition 8 *The following cases describe the conditions for always or never hiring the agent being optimal.*

- [1.] *If $\min\{1 - \bar{\pi} - c\Psi, 1 - \bar{\pi} + a\eta - c\} \geq 0$, $\chi(x) = 1$ for all $x \in \mathbb{R}_+ \cup \{\infty\}$ is optimal. The value function is given by $V^*(x) = 1 - \bar{\pi} + \frac{\Psi x}{1 + \Psi x} a\eta - \frac{1 + x}{1 + \Psi x} c\Psi$. If $a\eta > (1 - \Psi)c$, it is strictly increasing and strictly concave; if $a\eta < (1 - \Psi)c$, it is strictly decreasing and strictly convex. If $a\eta = (1 - \Psi)c$, $V^*(x) = 1 - \bar{\pi} - c\Psi$.*
- [2.] *If $\max\{1 - \bar{\pi} - c\Psi, 1 - \bar{\pi} + a\eta - c\} \leq 0$, $\chi(x) = 0$ for all $x \in \mathbb{R}_+ \cup \{\infty\}$ is optimal. The value function is $V^* = 0$ in this case.*

Proofs for our results rely on standard verification arguments; please see below for details.

In the following propositions, we shall show that, in the cases not covered by Proposition 8, the principal's learning benefit will be strictly positive,

¹⁰As we note in the main text, there is a discontinuity in payoffs at $x = 0$, which stems from the fact that, at $x = 0$, the contract we are looking at (payments contingent on success) ceases to be possible. As our contract continues to be possible, and (weakly) optimal, when $p^A = p^P = 1$, there is no such discontinuity at $x = \infty$.

and that her hiring decision will admit of a simple cutoff structure. First, if $1 - \bar{\pi} - c\Psi < 0 < 1 - \bar{\pi} + a\eta - c$, i.e., if the extra profit is important to the principal, meaning that η is large, and the initial disagreement regarding the agent's talent is not too severe, i.e., Ψ is not too low, the principal will hire the agent if and only if he is optimistic enough about his talent, as the following proposition shows.

Proposition 9 *If $1 - \bar{\pi} - c\Psi < 0 < 1 - \bar{\pi} + a\eta - c$, $\chi = \mathbb{1}_{(x^*, \infty]}$, with $x^* = \frac{r}{r+a}x^m$, is optimal. The value function is C^1 and given by*

$$V^*(x) = \mathbb{1}_{(x^*, \infty]}(x) \left[\frac{x^{-\frac{r}{a}}C}{1 + \Psi x} + 1 - \bar{\pi} + \frac{\Psi x}{1 + \Psi x}a\eta - \frac{1 + x}{1 + \Psi x}c\Psi \right],$$

where $C = -x^{*\frac{r}{a}}(1 + \Psi x^*) \left[1 - \bar{\pi} + \frac{\Psi x^*}{1 + \Psi x^*}a\eta - \frac{1 + x^*}{1 + \Psi x^*}c\Psi \right]$ is a constant of integration determined by value matching at $x = x^*$. On (x^*, ∞) , V^* is strictly increasing, and strictly convex (concave) on $(x^*, \tilde{x})$ ($(\tilde{x}, \infty)$), for some inflection point $\tilde{x} \in (x^*, \infty)$.

In this case, the principal will either never hire the agent if $x_0 \leq x^*$, or, if $x_0 > x^*$, he will initially hire the agent and keep hiring him until the time τ at which the belief $x_\tau = x^*$; the agent is fired for good at this time τ . The firing time $\tau = \tau^*$, where $\tau^* := \frac{1}{a} \ln(x_0/x^*)$, if the agent produces no extra profit η in $[0, \tau^*]$; otherwise, $\tau = \infty$, i.e., the agent is hired forever. This case is equivalent to a standard one-armed Poisson bandit problem, in which the risky arm is pulled whenever the decision maker is optimistic enough about its quality. The value function in this case is smooth, verifying the usual *smooth pasting* property. In our case, a success is fully revealing, so that the risky arm will be used forever after a success. In the absence of a success, optimism about its quality wanes continuously; the risky arm will be abandoned forever when beliefs hit a threshold (or we start out below this threshold). The principal's learning benefit shows up in the fact that she will hire the agent below the myopic cutoff x^m ; indeed, on $\left(\frac{1}{1+\frac{a}{r}}x^m, x^m\right)$, she is hiring the agent, even though her current payoffs would be higher if she produced herself. The concept of forgoing current payoffs in exchange

for information that is then parlayed into better decisions in the future is what the literature commonly refers to as *experimentation*. The extent of experimentation in our model is governed by the discounted arrival rate of information $\frac{a}{r}$; it vanishes as the principal becomes myopic ($r \rightarrow \infty$), and becomes large as information arrives quickly (a large).

If, however, $1 - \bar{\pi} + a\eta - c < 0 < 1 - \bar{\pi} - c\Psi$, i.e., if η and Ψ are relatively small, the opposite dynamics obtain. In this case, the extra profit is relatively unimportant to the principal, and the initial disagreement concerning the agent's talent is large. Then, the principal will hire the agent if and only if he is *pessimistic* enough about his talent, as the following proposition details.

Proposition 10 *If $1 - \bar{\pi} + a\eta - c < 0 < 1 - \bar{\pi} - c\Psi$, $\chi = \mathbb{1}_{[0, \check{x}]}$, with $\check{x} = \frac{a+r}{r}x^m$, is optimal. The value function in this case is given by $V^*(x) = \mathbb{1}_{[0, \check{x}]}(x) \left[1 - \bar{\pi} + \frac{\Psi x}{1 + \Psi x} a\eta - \frac{1+x}{1 + \Psi x} c\Psi - \frac{a}{a+r} \frac{\Psi x}{1 + \Psi x} (1 - \bar{\pi} + a\eta - c) \right]$; it is C^1 , except for a convex kink at $\check{x}$, flat on $[\check{x}, \infty)$, and strictly decreasing and strictly convex on $(0, \check{x})$.*

In this case, the principal will either never hire the agent if $x_0 > \check{x}$, or, if $x_0 \leq \check{x}$, she will hire the agent until he produces the extra profit, at which time she will fire him forever. In this case, the stopping boundary is not a *regular* boundary, as beliefs can only move away from, rather than toward, the boundary $\check{x}$, over the course of time. As in Keller and Rady (2015), therefore, *smooth pasting* fails, and the value function admits a kink at the boundary. As in the previous case, the extent of experimentation is increasing in the ratio $\frac{a}{r}$, with $\check{x} = \left(\frac{a}{r} + 1\right)^2 x^* = \left(\frac{a}{r} + 1\right) x^m$.

6.1 Proofs

6.1.1 Proof of Remark 1

We have to show that the principal cannot do better by ever paying the agent in the absence of a success. Suppose to the contrary that there exists a period t and a history such that the principal pays a flow $\hat{W}_t > 0$ in the

absence of a success and a lump sum of $W_t \geq 0$ after a success. Then, since at an optimum, the agent's participation constraint will bind, we have

$$\frac{x_t a}{1 + x_t} W_t + \hat{W}_t = c,$$

while the instantaneous (principal-)expected cost is

$$\frac{\Psi x_t a}{1 + \Psi x_t} W_t + \hat{W}_t.$$

Plugging in the agent's binding participation constraint yields

$$c - x_t a W_t \frac{1 - \Psi}{(1 + \Psi x_t)(1 + x_t)}.$$

As the factor multiplying $x_t a W_t$ is (strictly) less than 1 (if $\Psi < 1$), the principal has no incentive (a strict disincentive) to set $\hat{W}_t > 0$ (on a set of histories with positive measure, if $\Psi < 1$). Thus, it is optimal for the principal to set $W_t = \frac{1+x_t}{ax_t} c$ (a.s.), leading to a principal-expected cost of hiring of

$$\frac{\Psi x_t a}{1 + \Psi x_t} W_t = \frac{1 + x_t}{1 + \Psi x_t} \Psi c.$$

6.1.2 Proof of Lemma 1

The only claim that is not immediately obvious from inspection is that $\frac{1+x_t}{1+\Psi x_t} \Psi c$ is a martingale on the principal's information filtration. We have

$$\begin{aligned} \mathbb{E} \left[d \frac{1+x}{1+\Psi x} \Psi c \right] &= \frac{x \Psi a}{1 + \Psi x} c dt + \left(1 - \frac{x \Psi a}{1 + \Psi x} dt \right) \left[\frac{1+x}{1+\Psi x} \Psi c - x a \frac{1-\Psi}{1+\Psi x} \right] \\ &= \frac{\Psi c}{1 + \Psi x} dt \left\{ x a - \frac{x \Psi a}{1 + \Psi x} (1+x) - x a \frac{1-\Psi}{1 + \Psi x} \right\} + o(dt) = o(dt). \end{aligned}$$

6.1.3 Proof of Propositions 8–10

We shall write

$$\hat{V}(x) = 1 - \bar{\pi} + \frac{\Psi x}{1 + \Psi x} a\eta - c\Psi \frac{1 + x}{1 + \Psi x} - \mathbb{1}_{\{1 - \bar{\pi} + a\eta - c < 0\}} \frac{a}{a + r} \frac{\Psi x}{1 + \Psi x} (1 - \bar{\pi} + a\eta - c)$$

for the principal's payoff of never firing the agent in the absence of a success.

In all four cases, the proposed policy χ implies a well-defined law of motion of the belief x , and the closed-form expression for V^* is the payoff function associated with the policy χ . To prove optimality of χ , it suffices to show that $\mathcal{B}(x, V^*) \geq -\mathcal{M}(x)$ ($\mathcal{B}(x, V^*) \leq -\mathcal{M}(x)$) whenever $\chi = 1$ ($\chi = 0$) on some open subset of $\mathbb{R}_+$.

For Proposition 8, Case (1.), direct computation shows that $\mathcal{B}(x, \hat{V}) \geq -\mathcal{M}(x)$ for all $x \geq 0$. Moreover, $\hat{V}' > 0 > \hat{V}''$ if $a\eta > (1 - \Psi)c$, $\hat{V}' < 0 < \hat{V}''$ if $a\eta < (1 - \Psi)c$, and $\hat{V} = 1 - \bar{\pi} - c\Psi$ if $a\eta = (1 - \Psi)c$.

In Case (2.), $\mathcal{B}(x, V^*) = \mathcal{B}(x, 0) = 0$, for all $x \geq 0$. Thus, all that remains to be shown is that $\mathcal{M}^*(x) \leq 0$ for all $x \geq 0$. As $\mathcal{M}$ is increasing, this is equivalent to $\lim_{x \rightarrow \infty} \mathcal{M}(x) = 1 - \bar{\pi} - c + a\eta \leq 0$, which holds by the definition of Case (2.).

Let us turn to Proposition 9. For $x < x^*$, $V^*(x) = 0$ and $\mathcal{B}(x, V^*) = \frac{\Psi x a}{r(1 + \Psi x)} (1 - \bar{\pi} + a\eta - c)$. Direct computation shows that $\mathcal{B}(x, V^*) \leq -\mathcal{M}(x)$ for $x < x^*$. For $x > x^*$, one shows by direct computation that $\mathcal{B}(\cdot, V^*) > -\mathcal{M}(\cdot)$ in this range. Thus, $\chi = \mathbb{1}_{(x^*, \infty]}$ is optimal. Direct computation furthermore shows that $\lim_{x \downarrow x^*} V^{*'}(x) = 0$ and $V^{*'}(x) > 0$ for all $x > x^*$. By the same token, direct computation shows that $\lim_{x \downarrow x^*} V^{*''}(x) > 0$, $\lim_{x \rightarrow \infty} V^{*''}(x) < 0$, while $V^{*'''}|_{(x^*, \infty)} < 0$.

We now turn to Proposition 10. For $x > \check{x}$, $V^*(x) = \mathcal{B}(x, V^*) = 0$. By the same token, $\mathcal{M}(x) \leq 0$ if and only if $x \geq x^m = \frac{r}{a+r} \check{x}$. For $x < \check{x}$, one shows by direct computation that $\mathcal{B}(\cdot, V^*) > -\mathcal{M}(\cdot)$ in this range. Thus, $\chi = \mathbb{1}_{(0, \check{x}]}$ is optimal. Direct computation furthermore shows that $V^{*''}|_{(0, \check{x}]} > 0$, and that $\lim_{x \uparrow \check{x}} V^{*'}(x) < 0$.

6.1.4 Proof of Proposition 4

Note that $\lim_{W \rightarrow 0} u'(W) = \infty$ implies that w is positive for all x . Assuming that the principal hires the agent, maximizing the principal's myopic profits $(1 + p^P a(\eta - b) - w) dt$ subject to (PC) – which clearly binds in a profit-maximizing equilibrium – yields the following Lagrangian and first-order conditions

$$\begin{aligned}
L &= 1 + \frac{\Psi x}{(1 + \Psi x)} a(\eta - b) - w \\
&\quad + \lambda_{PC} \left[a \frac{x}{(1 + x)} u(w + b) + \left(1 - a \frac{x}{(1 + x)} \right) u(w) - c \right] \\
\frac{\partial L}{\partial w} &= -1 + \lambda_{PC} \left[a \frac{x}{(1 + x)} u'(w + b) + \left(1 - a \frac{x}{(1 + x)} \right) u'(w) \right] = 0 \\
\Rightarrow \lambda_{PC} &= \frac{1}{\left[a \frac{x}{(1 + x)} u'(w + b) + \left(1 - a \frac{x}{(1 + x)} \right) u'(w) \right]} \\
\frac{\partial L}{\partial b} &= -\frac{\Psi x}{(1 + \Psi x)} a + \lambda_{PC} \left[a \frac{x}{(1 + x)} u'(w + b) \right] = 0 \\
\Rightarrow -\frac{\Psi x}{(1 + \Psi x)} a + \frac{\left[a \frac{x}{(1 + x)} u'(w + b) \right]}{\left[a \frac{x}{(1 + x)} u'(w + b) + \left(1 - a \frac{x}{(1 + x)} \right) u'(w) \right]} &= 0 \\
\Rightarrow u'(w + b) [1 + (1 - a) \Psi x] - \Psi [1 + x(1 - a)] u'(w) &= 0
\end{aligned}$$

Note that also the sufficient condition for a maximum holds, therefore optimality conditions are

$$[1 + (1 - a) \Psi x] u'(w + b) - \Psi [1 + (1 - a) x] u'(w) = 0 \quad (\text{FOC})$$

$$axu(w + b) + [1 + (1 - a) x] u(w) - (1 + x) c = 0 \quad (\text{PC})$$

Since $\Psi [1 + (1 - a) x] / [1 + (1 - a) \Psi x] < 1$ for $\Psi < 1$, (FOC) implies that

$b > 0$ for all x , thus the agent is paid a success-based bonus irrespective of the extent of his risk aversion.

To show that W decreases in x (i.e., increases over time as long as no success is generated) with $u(W) = \ln W$, note that

$$\begin{aligned}
W &= w + b \\
&= w \frac{1 + \Psi(1-a)x}{\Psi[1 + (1-a)x]} \\
&\text{with} \\
\frac{\partial W}{\partial x} &= \frac{\partial w}{\partial x} \frac{1 + \Psi(1-a)x}{\Psi[1 + (1-a)x]} \\
&\quad - w \frac{(1-a)(1-\Psi)}{\Psi[1 + (1-a)x]^2} \\
&= - \frac{w \left[(1-a)(1-\Psi) + \frac{a[1 + \Psi(1-a)x] \ln \left(\frac{1 + (1-a)\Psi x}{\Psi[1 + (1-a)x]} \right)}{(1+x)} \right]}{\Psi(1+x)[1 + (1-a)x]} \\
&< 0
\end{aligned}$$

6.1.5 Proof of Proposition 5

We derive equilibrium contracts that maximize the agent's expected utility (according to the agent's assessment) subject to the constraint that the principal achieve an expected profit of at least 0, according to the principal's assessment. The principal's binding participation constraint pins down the agent's reward in case of a success $W_t = \frac{1 + \Psi x_t}{\Psi x_t a} (1 - \bar{\pi}) + \eta$, leading to an agent-expected myopic payoff for the agent of $\mathcal{M}_A(x) = \frac{1 + \Psi x}{\Psi(1+x)} (1 - \bar{\pi}) + \frac{xa}{1+x} \eta - c$. We note that this myopic payoff $\mathcal{M}_A$ is increasing (decreasing) if and only if $\Psi a \eta \geq (\leq) (1 - \Psi)(1 - \bar{\pi})$.

This in turn leads to a simple ODE for the agent's payoff $U(x)$,

$$x(1+x)aU'(x) + (r(1+x) + xa)U(x) = r(1+x)\mathcal{M}_A(x) + xa \max\{1 - \bar{\pi} + a\eta - c, 0\}.$$

Solving the ODE, and going through the same verification steps as in the baseline model¹¹ yields the exact same results as in the baseline model; i.e., the parameter ranges and threshold values of Propositions 1–3 continue to apply. We can thus conclude that both extreme allocations of bargaining power between the principal and the agent lead to the exact same results.

6.1.6 Proof of Proposition 7

The claim immediately follows from continuity and the fact that V_i^* is strictly decreasing in Ψ_i (when $\Psi_i \cdot x_i$ is held constant).

6.2 Microfoundation for Second Job

The purpose of this appendix is to show how to extend the model so as explicitly to incorporate the second job. Specifically, we shall denote $x_0 \in (0, \infty)$ ($\Psi_x x_0$) the agent's (principal's) belief (measured in odds ratios, as before) that the agent is talented for the first job, and hence produces the extra profit $\eta_x > 0$ at the rate $a_x > 0$ in the first job. By the same token, we shall write $y_0 \in (0, \infty)$ ($\Psi_y y_0$) for the agent's (principal's) belief that the agent is talented for the second job, and hence produces the extra profit $\eta_y > 0$ at the rate $a_y > 0$ in the second job. Flow opportunity costs in either job are $c_x > 0$, and $c_y > 0$, respectively.

We continue to assume that the agent is (weakly) overconfident regarding both jobs, i.e., that $\Psi_x \leq 1$ and $\Psi_y \leq 1$. Since talent across jobs is uncorrelated, we have $y_t = y_0$ for all times t at which the agent is employed in the first job. Both parties discount future payoffs at the rate $r > 0$. After the agent has been reassigned/promoted to the second job, the principal, as before, receives a flow payoff of $\bar{\pi}_y \geq 0$ if she does not hire the agent. Before the agent is reassigned, the principal receives a flow payoff of $\bar{\pi}_x \geq 0$ if she does not hire the agent. We shall write V_x^* for the agent's

¹¹Details are available from the authors upon request.

value to the principal in the first job, ignoring the possibility of reassignment to the second job. Clearly, the principal will reassign the agent at time $\tau^* = \inf \{t \geq 0 : \bar{\pi}_x + V_x^*(x_t) < \bar{\pi}_y + V_y^*(y_0)\}$.

The value functions V_x^* and V_y^* are computed as above. Before the agent is reassigned, y_t , and therefore $V_y^*(y_t) \equiv V_y^*(y_0)$, remain constant, while x_t , and hence $V_x^*(x_t)$, evolve as described above. The key to our subsequent analysis is the monotonicity of the value function, which we have noted in Remark 3. In particular V_i^* ($i \in \{x, y\}$) is strictly increasing (decreasing) if and only if $a_i\eta_i > (1 - \Psi_i)c_i$ ($a_i\eta_i < (1 - \Psi_i)c_i$), and constant if and only if $a_i\eta_i = (1 - \Psi_i)c_i$.

As before a reassignment, y_t , and hence $V_y^*(y_t)$, remain constant, only the monotonicity of V_x^* , and hence the properties of the first job, matter for the dynamics. In particular, for arbitrary parameters for the second job:

- If $a_x\eta_x > (1 - \Psi_x)c_x$, the agent is reassigned after a long enough dearth of lump sums $[0, \tau^*]$, with $\tau^* \in [0, \infty]$;
- if $a_x\eta_x < (1 - \Psi_x)c_x$, the agent is reassigned either right away, never, or at the arrival time of the first lump sum in the first job;
- if $a_x\eta_x = (1 - \Psi_x)c_x$, the agent is either reassigned right away or never.¹²

Reassignment dynamics thus depend only on the characteristics of the first job. In particular, the agent is reassigned after a long enough streak of failures if $a_x\eta_x > (1 - \Psi_x)c_x$. If $a_x\eta_x = (1 - \Psi_x)c_x$, his performance in the first job does not matter; he either stays in the first job forever, or is immediately affected to the second job. If $a_x\eta_x < (1 - \Psi_x)c_x$, the agent is reassigned/promoted as soon as he has proven his productivity in the first job by a success, which we interpret as a manifestation of the Peter Principle.

Thus, if $a_x\eta_x \leq (1 - \Psi_x)c_x$, the agent is either promoted right away or never in the absence of a success. If $a_x\eta_x > (1 - \Psi_x)c_x$, however, the agent is never

¹²This is neglecting the knife-edge case where $V_x^* = 1 - \bar{\pi}_x - \Psi_x c_x = V_y^*(y_0)$; in this case, the principal is indifferent over all promotion times in $[0, \infty]$, independently of the history.

reassigned after a success, but, in the absence of a success, may be reassigned at any time $\tau^* \in [0, \infty]$, the exact realization of which depends on the precise parameter values.

References

- BENSON, A., D. LI, AND K. SHUE (2019): “Promotions and the Peter Principle,” *Quarterly Journal of Economics*, 134, 2085–2134, 10.1093/qje/qjz022. (document), 1, 1, 4, 4, 4.4
- BILLETT, M. T., AND Y. QIAN (2008): “Are Overconfident CEOs Born or Made? Evidence of Self-Attribution Bias from Frequent Acquirers,” *Management Science*, 54, 1037–1051, 10.1287/mnsc.1070.0830. 8
- BONDT, W. F. D., AND R. H. THALER (1995): “Chapter 13 Financial decision-making in markets and firms: A behavioral perspective,” in *Handbooks in Operations Research and Management Science*: Elsevier, 385–410, 10.1016/s0927-0507(05)80057-x. 1
- BONNEY, L., C. R. PLOUFFE, B. HOCHSTEIN, AND L. L. BEELER (2020): “Examining salesperson versus sales manager evaluation of customer opportunities: A psychological momentum perspective on optimism, confidence, and overconfidence,” *Industrial Marketing Management*, 88, 339–351, 10.1016/j.indmarman.2020.05.012. 4.4
- CARD, D. (2022): “Who Set Your Wage?” *American Economic Review*, 112, 1075–1090, 10.1257/aer.112.4.1075. 3.3.3
- DANIEL, K., D. HIRSHLEIFER, AND A. SUBRAHMANYAM (1998): “Investor Psychology and Security Market Under- and Overreactions,” *The Journal of Finance*, 53, 1839–1885, 10.1111/0022-1082.00077. 8
- DELLAVIGNA, S., AND U. MALMENDIER (2004): “Contract Design and Self-Control: Theory and Evidence,” *Quarterly Journal of Economics*, 119, 353–402, 10.1162/0033553041382111. 1

- DEVARO, J., AND M. WALDMAN (2012): “The Signaling Role of Promotions: Further Theory and Empirical Evidence,” *Journal of Labor Economics*, 30, 91–147, 10.1086/662072. 1
- DIXON, M., AND B. ADAMSON (2011): *The Challenger Sale: Taking Control of the Customer Conversation*: Penguin Publishing Group, <https://books.google.de/books?id=pioPC90iMdMC>. 4.4
- ELIAZ, K., AND R. SPIEGLER (2006): “Contracting with Diversely Naive Agents,” *Review of Economic Studies*, 73, 689–714, 10.1111/j.1467-937x.2006.00392.x. 2
- ENGLMAIER, F., M. FAHN, AND M. A. SCHWARZ (2020): “Long-Term Employment Relations when Agents are Present Biased,” Working Paper. 1
- FAIRBURN, J. A., AND J. M. MALCOMSON (2001): “Performance, Promotion, and the Peter Principle,” *Review of Economic Studies*, 68, 45–66, 10.1111/1467-937x.00159. 1
- GERVAIS, S., J. B. HEATON, AND T. ODEAN (2011): “Overconfidence, Compensation Contracts, and Capital Budgeting,” *Journal of Finance*, 66, 1735–1777, 10.1111/j.1540-6261.2011.01686.x. 7
- GITTINS, J. (1974): “A dynamic allocation index for the sequential design of experiments,” *Progress in statistics*. 6
- GLASER, M., T. LANGER, AND M. WEBER (2012): “True Overconfidence in Interval Estimates: Evidence Based on a New Measure of Miscalibration,” *Journal of Behavioral Decision Making*, 26, 405–417, 10.1002/bdm.1773. 3.3.1
- GOEL, A. M., AND A. V. THAKOR (2008): “Overconfidence, CEO Selection, and Corporate Governance,” *Journal of Finance*, 63, 2737–2784, 10.1111/j.1540-6261.2008.01412.x. 7

- GROSSMAN, Z., AND D. OWENS (2012): “An unlucky feeling: Overconfidence and noisy feedback,” *Journal of Economic Behavior and Organization*, 84, 510–524, 10.1016/j.jebo.2012.08.006. 1
- GRUBB, M. D. (2015): “Overconfident consumers in the marketplace,” *Journal of Economic Perspectives*, 29, 9–36. 2
- HEIDHUES, P., AND B. KŐSZEGI (2010): “Exploiting Naivete about Self-Control in the Credit Market,” *American Economic Review*, 100, 2279–2303. 1
- HEIDHUES, P., B. KŐSZEGI, AND P. STRACK (2018): “Unrealistic Expectations and Misguided Learning,” *Econometrica*, 86, 1159–1214, 10.3982/ecta14084. 1
- HEIDHUES, P., B. KOSZEGI, AND P. STRACK (2021): “Convergence in models of misspecified learning,” *Theoretical Economics*, 16, 73–99, 10.3982/te3558. 1
- HESTERMANN, N., AND Y. L. YAOUANQ (2021): “Experimentation with Self-Serving Attribution Biases,” *American Economic Journal: Microeconomics*, 13, 198–237, 10.1257/mic.20180326. 1
- HOFFMAN, M., AND S. V. BURKS (2020): “Worker overconfidence: Field evidence and implications for employee turnover and firm profits,” *Quantitative Economics*, 11, 315–348, 10.3982/qe834. 1
- HUFFMAN, D., C. RAYMOND, AND J. SHVETS (2022): “Persistent Overconfidence and Biased Memory: Evidence from Managers,” *American Economic Review*, 112, 3141–3175, 10.1257/aer.20190668. 1
- HUMPHERY-JENNER, M., L. L. LISIC, V. NANDA, AND S. D. SILVERI (2016): “Executive overconfidence and compensation structure,” *Journal of Financial Economics*, 119, 533–558, 10.1016/j.jfineco.2016.01.022. 1
- KAHN, L. B., AND F. LANGE (2014): “Employer Learning, Productivity, and the Earnings Distribution: Evidence from Performance Measures,”

- The Review of Economic Studies*, 81, 1575–1613, 10.1093/restud/rdu021. 1
- KELLER, G., AND S. RADY (2015): “Breakdowns,” *Theoretical Economics*, 10, 175–202, 10.3982/te1670. 6
- LANGE, F. (2007): “The Speed of Employer Learning,” *Journal of Labor Economics*, 25, 1–35, 10.1086/508730. 1
- LARKIN, I., L. PIERCE, AND F. GINO (2012): “The psychological costs of pay-for-performance: Implications for the strategic compensation of employees,” *Strategic Management Journal*, 33, 1194–1214, 10.1002/smj.1974. 1
- LAZEAR, E. P. (2004): “The Peter Principle: A Theory of Decline,” *Journal of Political Economy*, 112, S141–S163, 10.1086/379943. 1
- LEMIEUX, T., W. B. MACLEOD, AND D. PARENT (2009): “Performance Pay and Wage Inequality*,” *Quarterly Journal of Economics*, 124, 1–49, 10.1162/qjec.2009.124.1.1. 3.3.1
- MALMENDIER, U., AND G. TATE (2005): “CEO Overconfidence and Corporate Investment,” *Journal of Finance*, 60, 2661–2700, <http://www.jstor.org/stable/3694800>. 1, 7
- (2015): “Behavioral CEOs: The Role of Managerial Overconfidence,” *Journal of Economic Perspectives*, 29, 37–60, 10.1257/jep.29.4.37. 1, 7
- MANNING, A. (2021): “Monopsony in Labor Markets: A Review,” *ILR Review*, 74, 3–26, 10.1177/0019793920922499. 3.3.3
- MEIKLE, N. L., E. R. TENNEY, AND D. A. MOORE (2016): “Overconfidence at work: Does overconfidence survive the checks and balances of organizational life?” *Research in Organizational Behavior*, 36, 121–134, 10.1016/j.riob.2016.11.005. 1

- MENKHOFF, L., M. SCHMELING, AND U. SCHMIDT (2013): “Overconfidence, experience, and professionalism: An experimental study,” *Journal of Economic Behavior & Organization*, 86, 92–101, 10.1016/j.jebo.2012.12.022. 3.3.1
- MILGROM, P., AND J. ROBERTS (1988): “An Economic Approach to Influence Activities in Organizations,” *American Journal of Sociology*, 94, S154–S179, 10.1086/228945. 1
- MUROOKA, T., AND Y. YAMAMOTO (2021): “Misspecified Bayesian Learning by Strategic Players: First-Order Misspecification and Higher-Order Misspecification,” *Working Paper*. 1
- MYERS, D. G. (2010): *Social psychology*: McGraw-Hill, 610. 1
- OTTO, C. A. (2014): “CEO optimism and incentive compensation,” *Journal of Financial Economics*, 114, 366–404, 10.1016/j.jfineco.2014.06.006. 1
- PETER, L. J., AND R. HULL (1969): *The Peter Principle: Why Things Always Go Wrong*: Buccaneer Books. 4
- PRESMAN, E. L. (1991): “Poisson Version of the Two-Armed Bandit Problem with Discounting,” *Theory of Probability and Its Applications*, 35, 307–317, 10.1137/1135038. 6
- PUETZ, A., AND S. RUENZI (2011): “Overconfidence Among Professional Investors: Evidence from Mutual Fund Managers,” *Journal of Business Finance & Accounting*, 38, 684–712, 10.1111/j.1468-5957.2010.02237.x. 3.3.1
- ROBBINS, H. (1952): “Some aspects of the sequential design of experiments.” 6
- DE LA ROSA, L. E. (2011): “Overconfidence and moral hazard,” *Games and Economic Behavior*, 73, 429–451, 10.1016/j.geb.2011.04.001. 1, 1, 3.3.2
- SANTOS-PINTO, L. (2008): “Positive Self-image and Incentives in Organisations,” *Economic Journal*, 118, 1315–1332, 10.1111/j.1468-0297.2008.02171.x. 1, 1, 2, 3.3.2

- SANTOS-PINTO, L., AND L. E. DE LA ROSA (2020): “Overconfidence in labor markets,” *Handbook of Labor, Human Resources and Population Economics*, 1–42. 1
- SAUTMANN, A. (2013): “Contracts for Agents with Biased Beliefs: Some Theory and an Experiment,” *American Economic Journal: Microeconomics*, 5, 124–156, <http://www.jstor.org/stable/43189633>. 1, 1
- SCHUMACHER, H., AND H. C. THYSEN (2022): “Equilibrium contracts and boundedly rational expectations,” *Theoretical Economics*, 17, 371–414, 10.3982/te4231. 1
- SEVY, B. (2016): “Why Great Salespeople Make Terrible Sales Managers,” October, <https://www.forbes.com/sites/groupthink/2016/10/10/why-great-sales-people-make-terrible-sales-managers/>. 4.4
- THOMPSON, W. R. (1933): “On the Likelihood That one Unknown Probability Exceed Another in View of the Evidence of Two Samples,” *Biometrika*, 25, 285–294, 10.1093/biomet/25.3-4.285. 6
- WALDMAN, M. (1984): “Job Assignments, Signalling, and Efficiency,” *The RAND Journal of Economics*, 15, 255–267, 10.2307/2555679. 1
- YAOUANQ, Y. L., AND P. SCHWARDMANN (2022): “Learning About One’s Self,” *Journal of the European Economic Association*, 20, 1791–1828, 10.1093/jeea/jvac012. 1, 1, 3.1

document