

Koch, Alexander K.; Morgenstern, Albrecht; Raab, Philippe

Working Paper

An Experimental Test of Career Concerns

IZA Discussion Papers, No. 1405

Provided in Cooperation with:

IZA – Institute of Labor Economics

Suggested Citation: Koch, Alexander K.; Morgenstern, Albrecht; Raab, Philippe (2004) : An Experimental Test of Career Concerns, IZA Discussion Papers, No. 1405, Institute for the Study of Labor (IZA), Bonn

This Version is available at:

<https://hdl.handle.net/10419/20703>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.

IZA DP No. 1405

An Experimental Test of Career Concerns

Alexander K. Koch
Albrecht Morgenstern
Philippe Raab

November 2004

An Experimental Test of Career Concerns

Alexander K. Koch

University of London and IZA Bonn

Albrecht Morgenstern

Federal Ministry of Finance and IZA Bonn

Philippe Raab

*Bonn Graduate School of Economics
and IZA Bonn*

Discussion Paper No. 1405
November 2004

IZA

P.O. Box 7240
53072 Bonn
Germany

Phone: +49-228-3894-0
Fax: +49-228-3894-180
Email: iza@iza.org

Any opinions expressed here are those of the author(s) and not those of the institute. Research disseminated by IZA may include views on policy, but the institute itself takes no institutional policy positions.

The Institute for the Study of Labor (IZA) in Bonn is a local and virtual international research center and a place of communication between science, politics and business. IZA is an independent nonprofit company supported by Deutsche Post World Net. The center is associated with the University of Bonn and offers a stimulating research environment through its research networks, research support, and visitors and doctoral programs. IZA engages in (i) original and internationally competitive research in all fields of labor economics, (ii) development of policy concepts, and (iii) dissemination of research results and concepts to the interested public.

IZA Discussion Papers often represent preliminary work and are circulated to encourage discussion. Citation of such a paper should account for its provisional character. A revised version may be available directly from the author.

ABSTRACT

An Experimental Test of Career Concerns*

Holmström's (1982/99) career concerns model has become an important workhorse for the analysis of agency issues in many fields. The underlying signal jamming argument requires players to use information in a Bayesian way – which may or may not reasonably approximate real-life decision makers' behavior. Testing this theory with field data is difficult since typically little is known about the information that individuals base their decisions on, and this explains the dearth of empirical studies. We provide experimental evidence that the signal jamming mechanism works in a laboratory setting. Moreover, subjects' beliefs fit remarkably well requirements imposed by the Bayesian equilibrium concept.

JEL Classification: C91, D83, L14

Keywords: incentives, reputation, career concerns, signal jamming, experiments

Phillippe Raab
IZA Bonn
P.O. Box 7240
53072 Bonn
Germany
Email: raab@iza.org

* We are grateful to Armin Falk, Simon Gächter, Lars Koch, Hans Normann, and Jörg Oechssler for valuable comments. Moreover, we thank Bernd Irlenbusch and Dirk Sliwka for providing us with information on their experiments. Financial support from the Bonn Graduate School of Economics is gratefully acknowledged. This research was done while the first and second authors were at the University of Bonn. The opinions expressed in this paper are those of the authors and do not necessarily reflect the views of the Federal Ministry of Finance.

1 Introduction

Building on Holmström’s (1982/99) seminal work, the career concerns model has become an important workhorse for the analysis of agency issues. Both the model itself and the ‘signal jamming’ logic that it develops have been applied by economic theorists in a wide array of fields, including labor relationships and organizational economics (for a survey see Borland (1992)) as well as industrial organization (e.g., Fudenberg and Tirole (1986) and Samuelson, Mirman, and Schlee (1994)), financial economics (e.g., Stein (1989)), and political economy (e.g., Lohmann (1998, 1999) and LeBorgne and Lockwood (2003)). However, the abundance of theoretical work has been met by relatively few empirical studies on the effectiveness of career concerns (e.g., Gibbons and Murphy (1992) and Fee and Hadlock (2003)). A major difficulty in testing this theory with field data is that it is virtually impossible to control for the information that decision makers have at hand. In this paper we conduct a laboratory experiment to fill this gap and provide *direct* evidence on the empirical viability of the career concerns model.

In its most basic version, the career concerns argument goes as follows. An agent (e.g., a manager) can take an action (exert costly effort) to improve the signal (output) that principals (potential employers) receive about her productive abilities. Even though all parties are symmetrically informed, the agent has an incentive to interfere with the principals’ updating by exerting effort to make the signal revealed to them look more favorable (*signal jamming*). In equilibrium, principals take this into account and the agent’s signal is discounted accordingly. As a consequence, the agent is trapped into providing the effort to meet principals’ tougher standards and not to look worse than she actually is.

Our experimental design implements a simplified version of Holmström’s (1982/99) model. Managers are endowed with a random ability level which is unknown to them and the other subjects. In the first period they choose the level of costly effort. Together with the manager’s random ability level this effort determines the performance signal which is observed by the firm subjects. Firms are interested only in a manager’s ability and compete in the second period by making wage offers to managers. In the *Hidden Ability* treatment, which implements the simplified Holmström (1982/99) model, only the sum of each manager’s (unknown) ability and effort is revealed to the firm subjects before they make wage offers. In the *Revealed Ability* treatment, which serves as a control, both ability and effort are revealed separately before second period wage offers are made. The theoretical prediction is that positive effort levels should only be observed in the Hidden Ability treatment. In that treatment ability is (symmetrically) unknown to the players so that managers can influence the firms’ updating about their ability upwards by exerting effort. For the sake of comparability, we chose a design and parameters that are comparable to the experiment by Irlenbusch and Sliwka (2004a), who investigate the impact of transparency on incentives in a more complex informational environment that allows for both signaling and signal jamming (see below).

The first contribution of our experiment is that it delivers the first clean test of whether the qualitative predictions on signal jamming from the career concerns model hold up in the lab. This provides evidence on how solid the foundations of theories are that build on career concerns being effective in creating reputational incentives. The advantage of our controlled laboratory environment over field data is that we can even move beyond this. Finding that subjects tend to play according to the predictions derived using the signal jamming logic does not actually tell us about the reasoning that led subjects to such behavior. Theoretical career concerns models rely on the Bayesian equilibrium concept, which puts high demands on the

rationality of players and requires sophisticated reasoning about other players' strategies (e.g., Fudenberg and Tirole (1991)). The second contribution of our experiment is that it elicits subjects' first and second order beliefs, and investigates whether players correctly conjecture each others' actions and beliefs, which is central to the concept of Bayesian equilibrium.

To preview, our main finding is that career concerns are actually effective in providing effort incentives. In line with the theoretical predictions, our experimental agents invest more costly effort when their ability is not perfectly revealed to principals compared to the Revealed Ability setting. Moreover, we obtain both direct and indirect evidence that subjects adjust their beliefs as postulated by the theory. First, experimental firms make roughly zero profits, which requires accurate estimates of managers' ability levels despite managers' attempts to jam the signal that inference is based upon. Second, the measured first and second order beliefs of subjects are remarkably consistent with each other and actual play in the game, which is a precondition for the concept of Bayesian equilibrium to produce a reasonable approximation of the outcome of individual behavior under incomplete information. Our results are in contrast to Irlenbusch and Sliwka (2004a), who report for their experiment that average effort is higher whenever effort and ability can be observed separately. A comparison of our setups suggests that these effects are due to the signaling opportunities arising from repeated gift exchange in their experimental labor market and do not contradict the relevance of career concerns as such. We discuss this in more detail in Section 5.

The paper is organized as follows. Section 2 introduces the simple variant of Holmström's (1982/99) career concerns model that we implement in the laboratory and derives the theoretical predictions. Section 3 presents the experimental design and reports procedures. In Section 4, we analyze the results. Section 5 compares our findings with those in Irlenbusch and Sliwka (2004a). Section 6 concludes.

2 Experimental setup with career concerns

2.1 Model

Setup. The central feature of Holmström's (1982/99) model is that the labor market determines a manager's wage based on an imperfect signal about her ability. Therefore, we will refer to our simplified version of this setup as the *Hidden Ability* model. Although career concerns occur in a number of situations and our experiment has a neutral frame, we illustrate the argument using Holmström's original interpretation of a manager-labor market relationship. A manager of uncertain ability first chooses an effort level which influences the signal observed by the labor market. More specifically, the manager is endowed with an ability level¹

$$a \in \{0, 1, 2, \dots, 19\}, \quad (1)$$

which is unobservable both to the manager and the principals. What is common knowledge is that a is uniformly distributed on the set $\{0, \dots, 19\}$. The manager chooses some effort level

$$e \in \{0, 1, 2, \dots, 19\} \quad (2)$$

¹We depart from Holmström's version of the model by imposing a finite support on the ability parameter a to be able to implement the model in a laboratory setting. For the sake of comparability of experimental results, our model and parameter choices follow Irlenbusch and Sliwka (2004a).

at a private cost $c(e)$. These costs are increasing in e (for details see Appendix A). Together with the manager's random ability level this effort determines the performance signal

$$y = a + e, \quad (3)$$

which is observed by the labor market. Upon receiving this information, $N \geq 2$ firms $i \in \{1, \dots, N\}$ active in that market offer the manager a wage $w_i(y)$. If the manager accepts an offer $w_i(y)$ she has a payoff of

$$u(e, w_i(y)) = w_i(y) - c(e). \quad (4)$$

If she declines all offers she receives her outside option, which is normalized to zero, and still has to bear the cost $c(e)$. The employing firm i receives the manager's ability less the wage it offered. Thus, the firm's profit is²

$$\pi_i(a, w_i(y)) = a - w_i(y). \quad (5)$$

Firms that do not succeed in hiring the manager make zero profits.

Equilibria. To construct a Perfect Bayesian Equilibrium with effort level e^* , we proceed by solving backward. Clearly, in the continuation game following firms' wage offers, the manager maximizes her payoff by accepting the highest offer $w_{max}(y) = \max\{w_i(y) | i = 1, \dots, N\}$. Since there are multiple firms which all attach the same value to the manager's services there is Bertrand-like competition if all firms hold the same beliefs. Thus, in such an equilibrium at least two firms offer a wage equal to the expected ability of the manager given market belief e^* and conditional on the signal y , i.e.

$$w_{max}(y) = E[a | y, e^*]. \quad (6)$$

Since $y = a + e$ and firms hold the belief that the manager exerted effort $e = e^*$, we have $w_{max}(y) = E[a | y, e^*] = y - e^* = a$ for realizations that are on the equilibrium path. Firms' beliefs about a when they observe performance signals $y \in \{0, \dots, e^* - 1, e^* + 20, \dots, 38\}$ that are off the equilibrium path are not pinned down by the equilibrium strategies. To be consistent with the model structure these beliefs only have to fulfill the two following conditions:

$$E[a | y, e^*] \leq \min\{y, 19\}. \quad (\text{C } 1)$$

$$E[a | y, e^*] \geq \max\{y - 19, 0\}. \quad (\text{C } 2)$$

For example, a manager with output 2 can at most have ability 2 (implying that zero effort was provided) and a manager with output 38 must have ability 19 (implying that the maximum possible effort level 19 was chosen).

Given firms' beliefs the manager chooses effort $e \in \{0, \dots, 19\}$ to maximize her future wage

²Note that the manager's first-period effort does not affect firm profits and is purely wasteful. An example for this is given in footnote 19.

minus the effort costs, $E[w_{max}(y)|e] - c(e)$. Note that a manager's effort e greater (less) than the market's belief e^* increases (decreases) the perceived ability of the manager for all y which can be observed under the equilibrium effort e^* . That is $\forall y \in \{e^*, \dots, e^* + 19\}$. Note that for signals y which are not in the equilibrium set this need not hold true because beliefs are only constrained by the two conditions (C 1) and (C 2). In equilibrium, e^* and the corresponding beliefs must be such that the manager has no incentive to deviate from e^* . That is,

$$e^* = \operatorname{argmax}_{e \in \{0, \dots, 19\}} \{E[w_{max}(y)|e] - c(e)\}. \quad (7)$$

Although this means that, in equilibrium, the manager cannot fool firms in their assessment of ability a , she still has no incentive to deviate to effort levels below e^* . The labor market expects exactly e^* and discounts the signal y accordingly, so that any lower effort level would make a manager look worse than she actually is.

Without further assumptions on the beliefs off the equilibrium path, it turns out that this leaves room for multiple equilibria. In particular, it can be shown that there exist beliefs which support pure strategy equilibrium efforts $e^* = 0, \dots, 12$ (see appendix B).³ A unique equilibrium effort level $e^* = 12$ can be achieved by imposing the following monotonicity restriction on beliefs:

$$E[a | y, e^*] = \begin{cases} 0 & \text{for } 0 \leq y < e^* \\ y - e^* & \text{for } e^* \leq y < e^* + 20 \\ 19 & \text{for } e^* + 20 \leq y \leq 38 \end{cases} \quad (MR)$$

This yields the following result:

Proposition 1 (Hidden Ability model)

If firms observe signal $y = a + e$, there are pure strategy Perfect Bayesian Equilibria under which

- *the manager chooses an effort level $e^* \in \{0, \dots, 12\}$*
- *at least two firms offer $w_{max}(y) = y - e^*$.*

Under belief restriction (MR), all Perfect Bayesian equilibria involve effort level $e^ = 12$.*

Hence, career concerns suffice to motivate the manager to exert costly effort if all parties only observe an imperfect measure of ability that the agent can manipulate.

In contrast, without uncertainty about ability, there is no reason for the manager to manipulate her performance signal by exerting costly effort since firms are only interested in ability. Specifically, in what we will refer to as the *Revealed Ability* model firms can not only observe the performance signal y but also the manager's ability a prior to their wage offers. Because this is a game of complete information the comparable solution concept to the one used in the Hidden Ability model is Subgame Perfect Nash equilibrium. Solving backward, we know that

³This multiplicity arises because the performance signal y can take values that are off the equilibrium path due to the finite support of the ability parameter a , and for such realizations beliefs are not pinned down by Bayes' rule. This is in contrast to Holmström (1982/99) (and other theoretical models), who derives a unique equilibrium effort level in a model where a has full support.

Bertrand competition leads to the best wage offers received by the manager

$$w_{max}(y, a) = a. \quad (8)$$

Since the manager's chosen wage is independent of effort, and effort is costly, all Subgame Perfect Nash equilibria involve effort $e^* = 0$. This is summarized in the following proposition.

Proposition 2 (Revealed Ability model)

If firms observe separately a and e , all Subgame Perfect Nash equilibria have the following properties:

- *the manager chooses an effort level $e^* = 0$*
- *at least two firms offer $w_{max}(y, a) = a$.*

Based on Propositions 1 and 2 evidence for the effectiveness of career concerns in an experimental setting can be found through a comparison of the Revealed and the Hidden Ability models. Accordingly, the treatments which implement them are labelled *Hidden Ability* (HA) treatment and *Revealed Ability* (RA) treatment, respectively.

2.2 Research questions

Effort. Our first goal is to test whether the introduction of career concerns has a positive impact on effort incentives. That is, we want to establish whether the uncertainty about the manager's ability in the Hidden Ability treatment induces her to exert costly effort to generate a more favorable performance signal, using the Revealed Ability treatment as a control. If we impose no restrictions on the shape of subjects' beliefs, our theoretical predictions in Propositions 1 and 2 lead to the following hypothesis:

Hypothesis 1

In the Hidden Ability (HA) treatment, effort is weakly higher than in the Revealed Ability (RA) treatment:

$$e^{HA} \geq e^{RA}. \quad (H1)$$

As noted before, to be able to implement the model in a laboratory setting it is necessary to impose a finite support on the ability parameter a . As a consequence, the theory yields multiple equilibria for the HA treatment. In particular, the observation $e^{HA} = e^{RA} = 0$ could mean that either

1. there are no incentives from career concerns, or that
2. there are incentives from career concerns but players in the HA treatment happened to coordinate on the equilibrium in which $e^* = 0$.

Nevertheless, even though for technical reasons it is difficult to come up with a design that is simple enough to implement in the laboratory and that yields a unique theoretical prediction which allows us to *falsify* career concerns more easily, in our experimental setup we are able to establish their validity whenever subjects in the HA treatment play a strictly larger effort than under the RA treatment:

Remark 1

If effort in the Hidden Ability treatment is strictly higher than in the Revealed Ability treatment, i.e. $e^{HA} > e^{RA}$, this is strong evidence in favor of the effectiveness of career concerns.

Hence, our first objective will be to establish whether $e^{HA} > e^{RA}$.

Beliefs. Proposition 1 yields a unique equilibrium prediction if one imposes the monotonicity restriction (MR) on firm's beliefs. Thus, the following hypothesis presents a *joint* test of effectiveness of career concerns and the validity of condition (MR):⁴

Hypothesis 2

$$e^{HA} = 12, \quad e^{RA} = 0. \quad (\text{H2})$$

As noted above, a failure to reject Hypothesis 2 does not speak to the issue as to whether career concerns are empirically valid or not. Nevertheless, the theoretical concept of Perfect Bayesian Equilibrium imposes certain requirements on players' beliefs in the game, which however are weaker than condition (MR). First, players' beliefs have to be consistent with the model structure. Second, beliefs are required to be consistent with equilibrium actions. To analyze whether this is a good approximation of behavior in our experiment, we elicit firms' beliefs and managers' second order beliefs in the HA treatment (see the discussion in Section 4.2).

Wage offers and profits. The incentives provided by career concerns are clearly indirect: a higher performance signal in the HA treatment is not valuable per se but only has a value for the manager because it is likely to lead to higher wage offers to her. This mechanism must break down whenever ability is revealed. Hence, the theory predicts that, in the HA treatment, wage offers should depend on the performance signal (i.e. ability and effort) while, in the RA treatment, only ability should matter.

Hypothesis 3

In the HA treatment, wage offers depend on the performance signal, i.e. equally on effort and ability, while in the RA treatment, wage offers only depend on ability.

The theory also has clear predictions about profits in the different treatments. In particular, in the HA treatment, the manager's average profit should be (weakly) lower than in the RA treatment because average wages are the same but the manager exerts costly effort to manipulate the performance signal in the former and none in the latter. Bertrand competition implies that average firm profits should be zero, regardless of the information structure. However, while the variance of firm profits should be positive in the HA treatment there should be zero variance in profits in the RA treatment where ability is revealed prior to wage offers. This yields the following hypothesis.

Hypothesis 4

In the RA treatment, managers earn weakly more than managers in the HA treatment. Firms in both treatments earn zero profits on average but the variance of firm profits in the HA treatment is larger than in the RA treatment.

⁴Irlenbusch and Sliwka (2004a) compare similar treatments to ours but only test such a joint hypothesis.

3 Procedures

The experiment was carried out in June 2004 in the Experimental Laboratory of the University of Bonn. The 84 subjects were recruited using ORSEE (Greiner 1998) and the experiment was run computer based with the experimental software z-tree (Fischbacher 1999). The participants (44 percent female) were mostly students (93 percent), of which roughly one third were Economics majors. Less than 20 percent of the subjects had followed a game theory course. In total, we conducted six sessions which provided us with six independent observations for each treatment.

In the beginning of each session, subjects were randomly assigned to separate cubicles in which they remained throughout the entire experiment. Instructions were distributed to the subjects and read aloud.⁵ Any questions concerning the experiment were only allowed on a bilateral basis. After reading the instructions, subjects had to fill out a questionnaire to make sure that the instructions were well understood. Questions required them to calculate simple examples on how actions determine the round profit for a manager and for a firm. The experiment on the computer was only started after all subjects had successfully answered all questions.

In both treatments, a session consisted of two separate groups with seven participants each. Within each group, three subjects were assigned the manager role (*A-players*) and four subjects were assigned the firm role (*B-players*). Roles were kept fixed during the whole experiment.

In each session, subjects completed 12 rounds of play. Each round consisted of four stages. In the beginning of the first stage, each manager was endowed with an ability parameter (*base value*) drawn from the set $\{0, 1, \dots, 18, 19\}$, where each value had equal probability.⁶ This ability parameter was not revealed to anyone. That is, neither the firms nor the managers themselves knew the ability parameter. Then, each manager chose her effort level (*number*) from the set $\{0, 1, 2, \dots, 19\}$ (effort costs are given in Appendix A). In the beginning of the second stage, firms in the HA treatment were shown the performance signal of each manager (the sum of the current round's base value and number, labelled *result*). In the RA treatment, firms observed both the performance signal and the ability parameter for each manager. We randomized the order of appearance of agents' performance signals or abilities on the screen in every round. Subjects were told explicitly that they could therefore not infer the identity of the other subjects from the position of their signal on the screen. Given the information on performance and/or ability, each of the firms now made wage (*transfer*) offers from the set $\{0.00, 0.01, \dots, 38.00\}$ to each of the three managers in its group. In the third stage, every manager was presented with her four wage offers and could decide whether to accept one of them or reject all offers. Again, we randomized the ordering of the wage offers in each round and communicated this to the subjects. Each round concluded with the fourth stage in which the round profits were displayed. The round profit of a manager was the accepted wage offer (or zero if all offers were rejected) minus the cost of her chosen effort level. The round profit of a firm equalled the sum of the ability parameters of all the managers who decided to accept the firm's respective wage offers, minus the wages paid out to these managers.

After the eighth round in the HA treatment, we elicited subjects' beliefs. Firms were asked to indicate for every possible realization of the performance signal $y \in \{0, 1, 2, \dots, 38\}$ what ability parameter they would expect such a manager to have on average. In a similar fashion,

⁵To avoid framing effects we used a neutral representation of the model. A translation of the instructions is given in Appendix C.

⁶The three sequences of ability parameters for managers were drawn before the experiments and kept fixed across all observations to facilitate comparisons within and between treatments.

Table 1: Summary statistics for effort choices of managers

	Hidden Ability (HA)	Revealed Ability (RA)
Average effort	8.13	4.19
Standard deviation	5.86	4.98
Zero effort played (percent)	16.67	46.30

managers were asked to report what belief they would expect from a firm on average for a given signal realization (second order beliefs). At the end of the experiment, the twelve round profits and an initial endowment were added up to the total profit in terms of experimental tokens. Finally, subjects were given a short questionnaire after play had been concluded. Since each firm could hire as many managers as it wanted to, and the number of firms exceeded that of managers, we expected firms to compete aggressively in wage offers, resulting in roughly zero round profits for firms. To finance their wages, each firm received an initial endowment of 180 experimental tokens in the first round. Since managers were expected to earn significant round profits, and we wanted firms and managers to earn more or less the same in the experiment, the managers' initial endowment was adjusted to 80 experimental tokens.⁷

At the end of the experiment subjects got paid their total profit converted into Euros at an exchange rate of 0.05 Euros per token. In the HA treatment, we additionally paid 3 Euros to compensate for the longer duration of HA sessions due to the belief elicitation after the eighth round. In total, HA sessions lasted for roughly 1 hour and 20 minutes with average earnings of 10.72 Euro. RA sessions lasted for about 1 hour and average earnings were 8.94 Euro.

4 Results

4.1 Effort choices

A first impression of the impact of uncertainty on effort incentives is provided by the summary statistics in Table 1. On average, managers in the HA treatment chose an effort level of 8.13, and therefore almost four effort units more than their counterparts in the RA treatment, who only provided average effort 4.19. This difference is not only considerable in absolute terms but highly significant as well (p-value 0.008, Wilcoxon-Mann-Whitney (WMW) test).⁸ Figure 1 reveals that this pattern is stable across periods – in particular, average effort in the HA treatment is higher than in the RA treatment in any given round. Hence, we cannot only confirm hypothesis 1 but, by Remark 1, have also established the relevance of career concerns. Average effort in the RA treatment differs from the theoretical prediction of $e^{RA} = 0$. Note however that decision errors cannot wash out, as they would with a predicted interior solution, and thus necessarily increase observed average effort above the predicted level. Learning does not appear to drive these errors since the incidence of zero effort does not increase markedly over rounds. However, psychological feedback from round payoffs was rather weak. In contrast

⁷The instructions mentioned initial endowments but levels were only disclosed to the subjects assigned the respective role via the current balance on the computer screen to preclude any impact of perceived inequalities in endowments.

⁸In the paper, we report results based on all periods. All our qualitative findings also obtain when using only the pre-questionnaire periods 1-8.

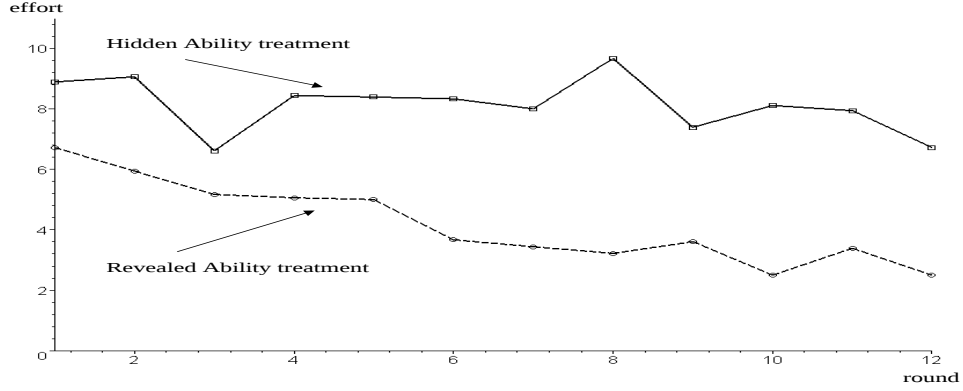


Figure 1: Average effort of managers across all rounds

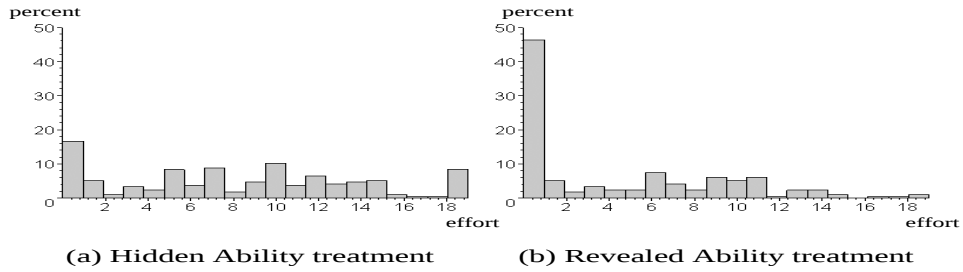


Figure 2: Effort choices of managers over all rounds

to firms, who realized a loss when offering too high wages, managers' effort decisions very rarely resulted in negative round payoffs. Nevertheless, managers provided zero effort in 46.30 percent of all cases, in accordance with the theoretical predictions. This is significantly more than in the HA treatment, where zero effort is provided in only 16.67 percent of the cases (p-value 0.021, WMW-test).⁹ A more detailed picture of the distribution of effort choices is given in Figure 2. In the RA treatment effort levels are concentrated on zero, whereas they are more dispersed in the HA treatment.

4.2 Beliefs

The results in the last section are strong evidence for the effectiveness of career concerns in providing effort incentives. A closer look at players' elicited beliefs allows us to address the issue of whether it is reasonable to impose the restriction (MR) on beliefs to achieve a unique

⁹This test compares the rankings of the number of times managers play zero in each group for each treatment. Note that a chi-squared test comparing the number of times zero effort is chosen relative to the number of times positive effort is chosen in the two treatments would require the assumption that effort decisions *within* each group were independent. If one is willing to make this assumption, the p-value is 0.001 (Chi-squared test).

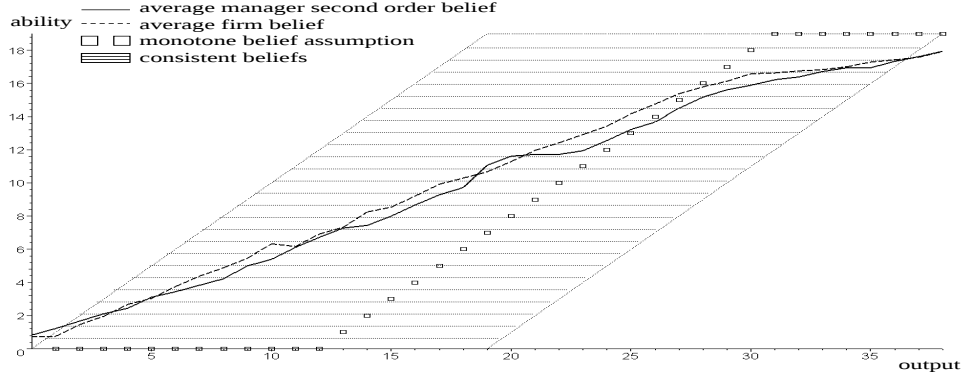


Figure 3: Average firm beliefs and average manager second order beliefs across all rounds

prediction under the theoretical model. Hypothesis 2 presents a joint test of the effectiveness of career concerns *and* the validity of condition (MR). However, the experimental data clearly contradict the point prediction that $e^{HA} = 12$ and $e^{HA} = 0$. Moreover, reported beliefs are very different from the ones posited by restriction (MR). This can be seen in Figure 3 by comparing the average first and second order beliefs with the profile of box symbols.

Nevertheless, a violation of (MR) does not imply that subjects behaved in a manner inconsistent with the theory. The first notable feature is that average reported beliefs (see Figure 3) are reasonable in the following sense. As discussed in Section 2.1, consistency with the model structure requires that beliefs must satisfy the restrictions (C 1) and (C 2). The hatched area marks the region where beliefs satisfy these consistency conditions. Even at an individual level beliefs are surprisingly consistent, out of the 1,575 reported beliefs less than 4 percent violate these conditions. Remarkably, the managers' second order beliefs correspond closely to the firms' beliefs. This can already be seen from the strong correlation of average values in Figure 3 and is also apparent at the individual level as shown by the boxplots¹⁰ of errors in managers' second order beliefs in Figure 4. These are computed by taking the difference between a manager's second order belief for a given output level and the average firm belief within the manager's group. The results confirm that second order beliefs are roughly unbiased and remarkably close to firm beliefs.

According to the theoretical prediction, beliefs should have a slope equal to unity on the equilibrium path, i.e. over a range of 19 different values of y . Reported beliefs typically do not display this feature because they are averages over the beliefs that a manager holds for three different managers within the same experimental group. This can be attributed to the strategic uncertainty arising from multiple equilibria in the game. Firm subjects appear to realize that managers play different effort levels¹¹ and thus discount higher realizations of the measure more than lower ones¹². Further support for this conjecture is provided by an

¹⁰The boxes span the range from first to third quartile of the data, with the median marked by a line in between. Whiskers extend to the most extreme data points within 3/2 of the interquartile range, beyond which extreme outliers are separately marked by circle symbols.

¹¹In the post experiment questionnaire 14 out of 24 firms report that they thought managers chose "very different effort levels" rather than "very similar effort levels."

¹²To see this point assume that a firm subject holds the belief that each manager plays according to a different 'equilibrium strategy', e.g. $e_1 = 0$, $e_2 = 6$, and $e_3 = 12$. Moreover, suppose that the firm has monotone beliefs. Then the belief about the average manager has a slope of 1/3 in the range $[0, 6]$, of 2/3 in the range $[6, 12]$, of 1

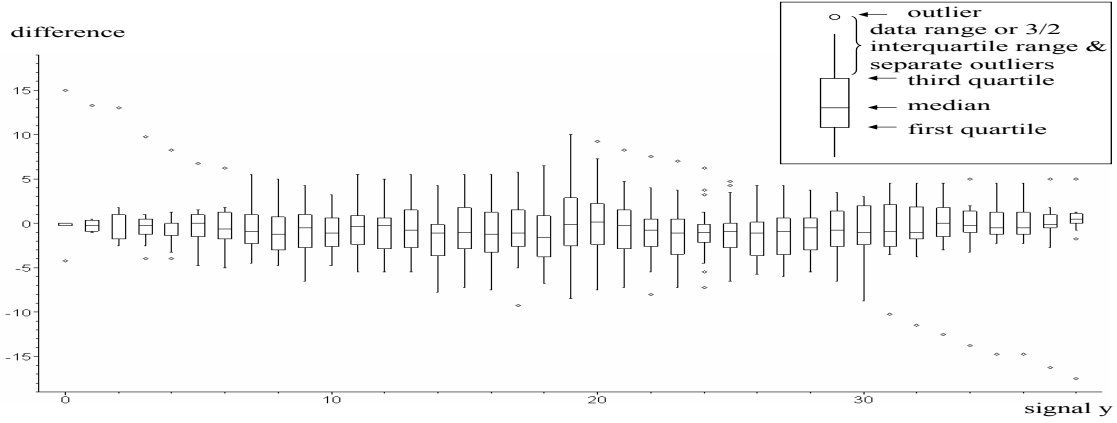


Figure 4: Errors in managers' second order beliefs across output levels

additional measure of firm beliefs that we obtained from the post experiment questionnaire in the HA treatment, where we ask for firm subjects' estimates of the average manager effort level. Strikingly, the firms learn quite well about managers' effort levels: the average reported belief is that managers provided effort 9.88 (standard deviation 3.35), which is not too far from the actual effort average level 8.13 (standard deviation 5.86) in the HA treatment. The remarkable degree of consistency in reported beliefs and the closeness of reported first and second order beliefs provide additional support for the hypothesis that, by and large, subjects have a good understanding of the signal jamming logic which provides them with effort incentives.

4.3 Wage offers and profits

A crucial ingredient in the theoretical model is that wage offers respond to effort levels – this we should see in the HA treatment, whereas effort should not impact wages in the RA treatment. To test this, we perform regressions of firms' wage offers on the respective manager's ability and effort levels in both treatments. As can be seen from Table 2 the experimental data nicely confirm our Hypothesis 3. In the HA treatment both effort and ability are significant. Moreover, the hypothesis that their coefficients are equal cannot be rejected (p-value: 0.81, Wald test) which is in line with the theoretical prediction. In contrast, in the RA treatment effort is not significant and wages react only to ability as predicted by the theory. A more fine-tuned prediction from the theory is that the coefficients on ability/effort (HA treatment) and ability (RA treatment) should be equal to one.¹³ The previous section discussed why in the HA treatment beliefs do not correspond to this prediction. The fact that the ability coefficient in the RA treatment is significantly lower than unity (p-value: 0.007, Wald test) partly reflects outlier behavior (see footnote 14). Nevertheless, Bertrand competition appears to work well in our setting as a look at average profits reveals.

Figure 5 plots average profits for managers and firms across rounds. It can clearly be seen that profits stabilize in both treatments after a few rounds. In contrast to hypothesis 4, managers make a somewhat higher profit in the HA treatment even though their effort costs are larger

in the range [12, 19], of 2/3 in the range [19, 25], of 1/3 in the range [25, 31], and is flat in the remaining part.

¹³Strictly speaking, Bertrand competition only predicts that this ought to hold for the highest two wage offers, but even then these coefficients are all less than one: the ability coefficients are 0.40 (HA) and 0.77 (RA) and the effort coefficients are 0.47 (HA) and -0.19 (RA).

Table 2: Wage regressions

	Hidden Ability (HA)	Revealed Ability (RA)
constant	1.49* (-0.71)	1.29 (0.77)
ability	0.37*** (0.05)	0.7*** (0.07)
effort	0.39*** (0.05)	-0.02 (0.04)
R ²	0.22	0.34
N	864	864

Numbers in parentheses are Huber-White robust standard errors, treating each experimental observation as a separate cluster.

Table 3: Summary statistics for manager and firm profits

	Hidden Ability (HA)	Revealed Ability (RA)
Average profit of managers	10.20	9.96
Standard deviation	8.74	8.35
Average profit of firms	-2.48	-1.40
Standard deviation	10.36	9.49

(see table 3). However, the differences are not significant (p-value: 0.409 (managers) and p-value: 0.155 (firms), WMW test) and the outcomes – particularly with respect to firms’ average profits – are broadly in line with the theory. That is, in the HA treatment initially firms incur substantial losses on average and then players learn to account for managers’ signal jamming to make roughly zero profits. As noted in Section 4.1, psychological feedback from profits was strong since firms that offered too high wages realized losses. In the RA treatment, profits very quickly converge to zero due to the competition for managers (even though the variance of profits is not much lower than for the HA treatment due to outliers).¹⁴ This nicely confirms the theoretical prediction of Bertrand competition in both treatments stated in Hypothesis 4 and is in line with other experimental evidence on Bertrand competition with more than two players (e.g. Dufwenberg and Gneezy (2000)).

¹⁴In each treatment there was one firm subject who went bankrupt. We provided them with credit and let them continue to play. Both subjects clearly misunderstood the experimental setup. The firm subject in the HA treatment who went bankrupt offered wages *larger* than the signal y that the manager produced in 40 percent of all his offers. He thus outbid all other firms and lowered average firm profits to -5.00 , compared to an average of -1.98 over the other HA observations. Moreover, his reported beliefs were not consistent with the experimental setup (see Section 4.2). The firm subject in the RA treatment who went bankrupt bid more than the *visible* ability parameter in 44 percent of his bids. This lowered the average firm profits to -4.26 compared to an average of -0.57 in the remaining RA observations.

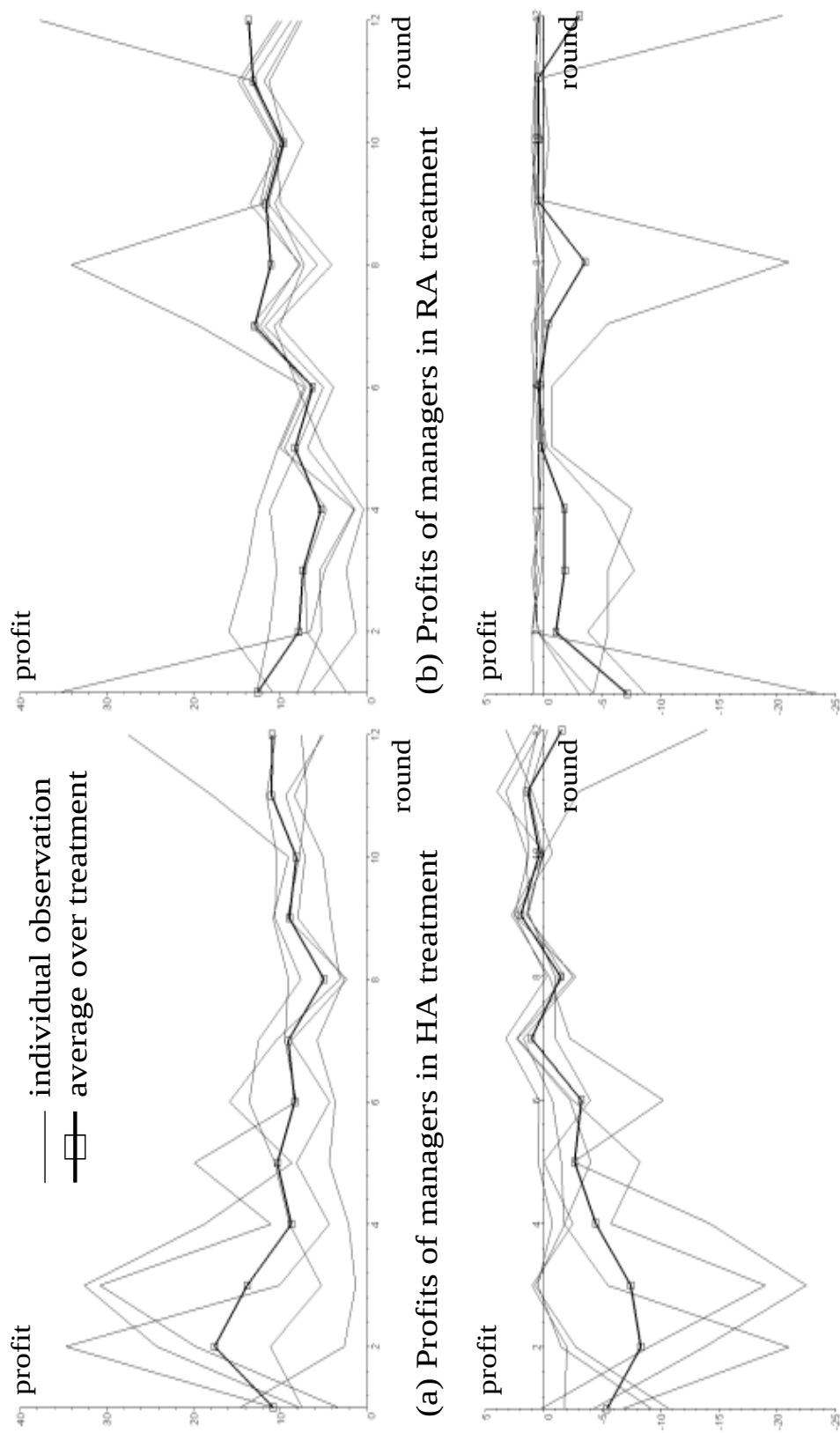


Figure 5: Profits of managers and firms across all rounds

Table 4: Comparison of results with Irlenbusch & Sliwka (2004a)

	Hidden Ability	Revealed Ability
	<i>Effort choices of managers</i>	
<i>our design</i>	8.13	4.19
<i>first period IS^a</i>	3.56	5.50
	<i>Coefficients in wage regressions^b</i>	
<i>our design</i>		
ability	0.37***	0.7***
effort	0.39***	-0.02
<i>first period IS^a</i>		
ability	{ 0.38**	0.841***
effort		0.373***

^a Irlenbusch & Sliwka (2004a)^b Regression of second period wage on first period ability/effort.

IS regressed the wage on output instead of ability/effort in their HA-treatment.

5 Career concerns and gift exchange

As noted before, we chose a design that is close to that of Irlenbusch and Sliwka (2004a) to be able to relate our findings to theirs. Our setup concentrates exclusively on the signal jamming element at the heart of career concerns models. In contrast, Irlenbusch and Sliwka (2004a) focus on the impact of transparency on incentives in the specific context of an experimental labor market as pioneered by Fehr, Kirchsteiger, and Riedl (1993), which leads to a more complex informational environment combining opportunities for both signaling and signal jamming.

Specifically, in their design firms set wages both in the beginning of the first and the second period (as opposed to only in the second period in our design).¹⁵ In both periods managers who accept a wage offer make an effort decision (our manager's only choose effort in the first period). The employing firm receives as payoff the sum of the manager's ability parameter and effort in the current period (whereas to our firms only the manager's ability matters).

For selfish preferences, the theoretical predictions for effort decisions in both setups are obviously identical. Table 4 summarizes experimental results. In Irlenbusch and Sliwka (2004a) first period average effort levels in the HA treatment are lower than those in the RA treatment. This is exactly the opposite of our finding. Differences in the dynamic structure of the two experiments in conjunction with other experimental evidence suggest that these findings are due to additional influences that interfere with the signal jamming mechanism in Irlenbusch and Sliwka (2004a).

In contrast to our design, firms in Irlenbusch and Sliwka (2004a) can 'invite' reciprocal behavior by making generous wage offers, and managers can reward such behavior through higher effort, which translates into more profit for the firm. It is well established from static versions of such experimental labor markets à la Fehr, Kirchsteiger, and Riedl (1993) that they encourage gift exchange in the form of more effort and higher wages than predicted by models based on selfish preferences. This finding is quite robust to the market institutions and their degree of

¹⁵More subtle differences to our design are that their setup allows transfers proposed by firms to managers to be either positive or negative and restricts them to the integer set $\{-38, \dots, 38\}$.

competitiveness.¹⁶

The dynamic structure introduces the possibility of additional influences. In Irlenbusch and Sliwka (2004a) subjects play a repeated gift exchange game (with the added feature that there is also learning about a manager’s unknown ability). If subjects perceive the experimental setup as a game of incomplete information about other players’ willingness to reciprocate, actions in the first period serve as signals about a player’s ‘type’. In such games, cooperative behavior can often be sustained for some time because players will engage in behavior in the first rounds to build a reputation for being a ‘cooperative’ type (e.g. Kreps, Milgrom, Roberts, and Wilson (1982)).¹⁷

Indeed, there is experimental evidence for such reputation building behavior in gift exchange games (Gächter and Falk (2002) and Irlenbusch and Sliwka (2004b)). The latter experiment implements a repeated gift exchange model similar to the one in Irlenbusch and Sliwka (2004a), replacing the permanent random ability component by a transitory productivity shock to output. In one treatment the firms can observe both profit components separately, while in the other treatment only the sum of effort and the productivity shock is revealed. The experiment shows that players’ reciprocal behavior is stronger, the less noisy the signal that the firm receives about the effort that a manager provided in response to the accepted wage. This ‘transparency effect’ on signaling about reciprocity is manifested in a significant increase in the wage-effort correlation when transparency on effort is increased.¹⁸

At first glance, one might expect that in the Irlenbusch and Sliwka (2004a) experiment the difference in effort levels across treatments is driven by two countervailing forces: the impact of transparency on career concerns (as witnessed in our experiment) and the transparency effect on signaling about reciprocity. Indeed, the higher average effort level in their RA treatment compared to our design supports this view (see Table 4). Moreover, first period effort has no impact on second period wages in our design, whereas the coefficient on effort is significant in the wage regression of Irlenbusch and Sliwka (2004a). This suggests that effort provides a means to signal the willingness to reciprocate in the second period – which is absent in our design. However, as Irlenbusch and Sliwka (2004a) note, *direct reciprocity*, as measured by the correlation between second period wages and second period effort, does not increase when moving from their HA to the RA treatment, in contrast to what could be expected from the results in Irlenbusch and Sliwka (2004b).

Another notable difference is that in our HA treatment average effort is higher than in Irlenbusch and Sliwka (2004a). This and the multiplicity of theoretically predicted signal jamming equilibria suggest that the effect from signaling about the degree to which a player is going to reciprocate in the second round cannot be neatly disentangled from the impact of signal jamming on effort induced by Holmström-type career concerns. As noted above, intransparency has a dampening effect on the use of effort as a signal of willingness to reciprocate in a repeated gift exchange setting. It is possible that this provides players with a means to coordinate on a low effort signal jamming equilibrium in the HA treatment. If players anticipate the correct effort level and have the right kind of beliefs, signal jamming could then not have any additional impact on effort beyond the one induced by signaling about reciprocity. This would offer

¹⁶E.g., Fehr, Kirchsteiger, and Riedl (1993), Fehr, Kirchsteiger, and Riedl (1998), Fehr, Kirchler, Weichbold, and Gächter (1998), Fehr and Falk (1999), and Brandts and Charness (2004).

¹⁷For a survey of experiments on reputation models and signaling in experiments more generally see Camerer (2003). Kübler, Normann, and Müller (2003) provide an experimental comparison of signaling and screening models.

¹⁸However, average first period effort levels do not change significantly when transparency is increased, only the dispersion of effort levels increases.

one explanation for the lower effort levels in the Irlenbusch and Sliwka (2004a) HA treatment compared to our experiment.

In sum, the comparison of the results from the two experiments shows that the findings of Irlenbusch and Sliwka (2004a) are overturned once one strips away the gift exchange components and allows only for signal jamming. This is an important insight. Clearly, the interaction of career concerns and gift exchange is relevant in many settings. Nevertheless, even in labor relations future career prospects may be very strongly influenced by the market's perception of an individual's ability (e.g., in the market for CEOs). In such circumstances, there is little reason to expect that gift exchange is an important factor in determining effort. The same is true whenever effort does not generate any additional surplus.¹⁹ Moreover, gift exchange is not relevant at all in many settings to which the career concerns model has been applied to²⁰. Overall, the results in our experiment and those in Irlenbusch and Sliwka (2004a),(2004b) suggest that the impact of increased transparency on effort depends very much on the nature of the interaction between the players. If gift exchange is an important element then transparency may induce more effort because of the possibility of using effort to signal being a reciprocator. In contrast, as our experiment demonstrates, if unknown ability matters a lot, transparency might reduce effort incentives because signal jamming is the dominant source of incentives for effort.

6 Conclusion

Models of signal jamming have become an important part of the toolkit of economic theorists. Despite the importance of this mechanism for theoretical work empirical evidence has been scant, due to limitations of field data. Our experiment provides a direct test of the signal jamming mechanism that forms the core of Holmström's (1982/99) model. The results support the theoretical premise that economic actors play according to the signal jamming logic. Moreover, subjects' beliefs that were elicited during the experiment are remarkably consistent with each other, which is a necessary condition for the concept of Bayesian equilibrium to produce a reasonable approximation of the outcome of individual behavior under incomplete information.

Our findings overturn those in the experiment by Irlenbusch and Sliwka (2004a), who analyze a setup where repeated gift exchange and career concerns interact. Overall, the evidence from the experimental studies suggest that the impact of transparency on incentives is sensitive to the exact nature of the interaction between the players. Transparency might enhance incentives whenever repeated gift exchange provides the possibility of using effort to signal willingness to reciprocate. In contrast, whenever actors have a strong interest in learning about unknown ability of other players transparency might reduce effort incentives created by the signal jamming mechanism.

¹⁹For example, consider a graduate student who prepares a job application. She may not know her ability (with respect to the tasks ahead) but her resumé provides a signal to potential employers. 'Polishing' the resumé is a costly activity that is likely to enhance the impression that potential employers get while it does not improve the applicant's ability or create any extra value to the future employer.

²⁰See the references listed in the introduction.

References

- Borland, Jeff, 1992, Career concerns: Incentives and endogenous learning in labour markets, Journal of Economic Surveys 6, 251–270.
- Brandts, Jordi, and Gary Charness, 2004, Do market conditions affect gift exchange? Evidence from experimental markets, Economic Journal 114, 684–708.
- Camerer, Colin F., 2003, Behavioral Game Theory (Princeton University Press: Princeton, N.J.).
- Dufwenberg, Martin, and Uri Gneezy, 2000, Price competition and market concentration: An experimental study, International Journal of Industrial Organization 18, 7–22.
- Fee, C. Edward, and Charles J. Hadlock, 2003, Raids, rewards and reputations in the market for CEO talent, Review of Financial Studies 16, 1315–1357.
- Fehr, Ernst, and Armin Falk, 1999, Wage rigidities in a competitive incomplete contract market. An experimental investigation, Journal of Political Economy 107, 106–134.
- Fehr, Ernst, Erich Kirchler, Andreas Weichbold, and Simon Gächter, 1998, When social norms overpower competition: Gift exchange in experimental labor markets, Journal of Labor Economics 16, 324–351.
- Fehr, Ernst, Georg Kirchsteiger, and Arno Riedl, 1993, Does fairness prevent market clearing? An experimental investigation, Quarterly Journal of Economics 108, 437–460.
- , 1998, Gift exchange and reciprocity in competitive experimental markets, European Economic Review 42, 1–34.
- Fischbacher, Urs, 1999, Z-tree - Zurich toolbox for readymade economic experiments - Experimenter’s manual, Working Paper Nr. 21, Institute for Empirical Research in Economics, University of Zurich.
- Fudenberg, Drew, and Jean Tirole, 1986, A “signal-jamming” theory of predation, Rand Journal of Economics 17, 366–376.
- , 1991, Game Theory (MIT Press: Cambridge, Mass.).
- Gächter, Simon, and Armin Falk, 2002, Reputation and reciprocity: Consequences for the labour relation, Scandinavian Journal of Economics 104, 1–27.
- Gibbons, Robert, and Kevin J. Murphy, 1992, Optimal incentive contracts in the presence of career concerns: Theory and evidence, Journal of Political Economy 100, 468–505.
- Greiner, Ben, 1998, An online recruitment system for economic experiments, Discussion Paper 10-2003, Max Planck Institute for Research into Economic Systems, Jena.
- Holmström, Bengt, 1982/99, Managerial incentive problems: A dynamic perspective, Review of Economic Studies 66, 169–182 ; originally published in: *Essays in Economics and Management in Honour of Lars Wahlbeck*, Helsinki, Finland.
- Irlenbusch, Bernd, and Dirk Sliwka, 2004a, Career concerns in a simple experimental labour market, European Economic Review forthcoming.

- , 2004b, Transparency and reciprocal behavior in employment relations, Journal of Economic Behavior and Organizations forthcoming.
- Kreps, David, Paul Milgrom, John Roberts, and Robert Wilson, 1982, Rational cooperation in the finitely repeated prisoners dilemma, Journal of Economic Theory 27, 245–252.
- Kübler, Dorothea, Hans Normann, and Wieland Müller, 2003, Job market signalling and screening in laboratory experiments, Royal Holloway, University of London.
- LeBorgne, Eric, and Ben Lockwood, 2003, Do elections always motivate incumbents? Experimentation vs. career concerns, IMF Working Paper 03/57.
- Lohmann, Susanne, 1998, Rationalizing the political business cycle: A workhorse model, Economics and Politics 10, 1–17.
- , 1999, What price accountability? The Lucas island model and the politics of monetary policy, American Journal of Political Science 43, 396–430.
- Samuelson, Larry, Leonard J. Mirman, and Edward E. Schlee, 1994, Strategic information transmission in duopolies, Journal of Economic Theory 62, 363–384.
- Stein, Jeremy C., 1989, Efficient capital markets, inefficient firms: A model of myopic corporate behavior, Quarterly Journal of Economics pp. 655–669.

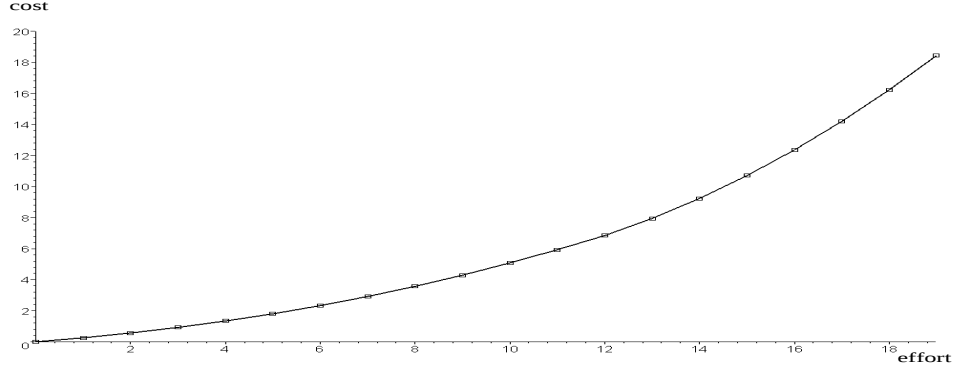


Figure 6: Cost function for effort.

Table 5: Cost function for effort.

	Effort ([e])									
	0	1	2	3	4	5	6	7	8	9
c(e)	0	0.26	0.57	0.93	1.34	1.8	2.33	2.92	3.57	4.29
	Effort ([e])									
	10	11	12	13	14	15	16	17	18	19
c(e)	5.08	5.93	6.84	7.94	9.23	10.7	12.35	14.19	16.22	18.43

A Cost function

B Derivation of equilibria

In this section we show how to construct Perfect Bayesian equilibria with a pure strategy equilibrium effort level $e^* \in \{0, \dots, 12\}$. Let $b[a|y, e^*]$ denote firms' beliefs about the manager's ability a , given that signal y has been observed and that the firms expect the manager to choose effort e^* . Table 6 provides examples of such beliefs supporting pure strategy equilibria. To illustrate the construction of an equilibrium, let us focus on the candidate equilibrium $e^* = 2$ (column four in Table 6). Suppose that a signal $y = 10$ realizes. This is on the (proposed) equilibrium path and thus $b[a|10, 2] = 10 - 2 = 8$. In contrast, $y = 1$ would be off the equilibrium path and beliefs only need to satisfy restrictions (C 1) and (C 2), as $b[a|1, 2] = 1$ does.

Given that firms hold beliefs $b[a|y, e^*]$, a manager chooses effort e to maximize the expected wage net of effort costs, $E[b[a|y, e^*]|e] - c(e)$, using the fact that a is uniformly distributed on $\{0, 1, 2, \dots, 19\}$. In our example, these net expected payoffs are given in column seven in Table 7. Column eight shows that the net expected payoff increases when moving from $e=0$ to $e=2$ and then always decreases, confirming that a manager will actually choose the conjectured equilibrium effort level $e^* = 2$. The derivations of the other equilibrium effort levels are given in Tables 7-11. The equilibrium with $e^* = 12$ in Table 11 corresponds to the unique equilibrium effort level if beliefs are restricted by assuming that the monotonicity condition (MR) holds. It is easy to verify that it is not possible to construct equilibria with effort levels $e^* > 12$.

$y = a + e$	equilibrium effort level (e^*)												
	0	1	2	3	4	5	6	7	8	9	10	11	12 ^a
	firm beliefs $b[a y, e^*]$ ^b												
0.0	0.0	<i>0.0</i>	<i>0.0</i>	<i>0.0</i>	<i>0.0</i>	<i>0.0</i>	<i>0.0</i>	<i>0.0</i>	<i>0.0</i>	<i>0.0</i>	<i>0.0</i>	<i>0.0</i>	<i>0.0</i>
1.0	1.0	0.0	<i>1.0</i>	<i>1.0</i>	<i>1.0</i>	<i>1.0</i>	<i>1.0</i>	<i>1.0</i>	<i>1.0</i>	<i>1.0</i>	<i>1.0</i>	<i>1.0</i>	<i>0.0</i>
2.0	2.0	1.0	0.0	<i>2.0</i>	<i>2.0</i>	<i>2.0</i>	<i>2.0</i>	<i>2.0</i>	<i>2.0</i>	<i>2.0</i>	<i>2.0</i>	<i>2.0</i>	<i>0.0</i>
3.0	3.0	2.0	1.0	0.0	<i>3.0</i>	<i>3.0</i>	<i>3.0</i>	<i>3.0</i>	<i>3.0</i>	<i>3.0</i>	<i>3.0</i>	<i>2.5</i>	<i>0.0</i>
4.0	4.0	3.0	2.0	1.0	0.0	<i>4.0</i>	<i>4.0</i>	<i>4.0</i>	<i>4.0</i>	<i>4.0</i>	<i>3.0</i>	<i>2.5</i>	<i>0.0</i>
5.0	5.0	4.0	3.0	2.0	1.0	0.0	<i>5.0</i>	<i>5.0</i>	<i>5.0</i>	<i>4.3</i>	<i>3.5</i>	<i>2.5</i>	<i>0.0</i>
6.0	6.0	5.0	4.0	3.0	2.0	1.0	0.0	<i>6.0</i>	<i>6.0</i>	<i>4.8</i>	<i>3.5</i>	<i>2.7</i>	<i>0.0</i>
7.0	7.0	6.0	5.0	4.0	3.0	2.0	1.0	0.0	<i>6.5</i>	<i>5.0</i>	<i>3.5</i>	<i>2.7</i>	<i>0.0</i>
8.0	8.0	7.0	6.0	5.0	4.0	3.0	2.0	1.0	0.0	<i>5.2</i>	<i>3.8</i>	<i>2.7</i>	<i>0.0</i>
9.0	9.0	8.0	7.0	6.0	5.0	4.0	3.0	2.0	1.0	0.0	<i>3.8</i>	<i>2.7</i>	<i>0.0</i>
10.0	10.0	9.0	8.0	7.0	6.0	5.0	4.0	3.0	2.0	1.0	0.0	<i>2.7</i>	<i>0.0</i>
11.0	11.0	10.0	9.0	8.0	7.0	6.0	5.0	4.0	3.0	2.0	1.0	0.0	<i>0.0</i>
12.0	12.0	11.0	10.0	9.0	8.0	7.0	6.0	5.0	4.0	3.0	2.0	1.0	0.0
13.0	13.0	12.0	11.0	10.0	9.0	8.0	7.0	6.0	5.0	4.0	3.0	2.0	1.0
14.0	14.0	13.0	12.0	11.0	10.0	9.0	8.0	7.0	6.0	5.0	4.0	3.0	2.0
15.0	15.0	14.0	13.0	12.0	11.0	10.0	9.0	8.0	7.0	6.0	5.0	4.0	3.0
16.0	16.0	15.0	14.0	13.0	12.0	11.0	10.0	9.0	8.0	7.0	6.0	5.0	4.0
17.0	17.0	16.0	15.0	14.0	13.0	12.0	11.0	10.0	9.0	8.0	7.0	6.0	5.0
18.0	18.0	17.0	16.0	15.0	14.0	13.0	12.0	11.0	10.0	9.0	8.0	7.0	6.0
19.0	19.0	18.0	17.0	16.0	15.0	14.0	13.0	12.0	11.0	10.0	9.0	8.0	7.0
20.0	<i>4.8</i>	19.0	18.0	17.0	16.0	15.0	14.0	13.0	12.0	11.0	10.0	9.0	8.0
21.0	<i>6.7</i>	<i>5.7</i>	19.0	18.0	17.0	16.0	15.0	14.0	13.0	12.0	11.0	10.0	9.0
22.0	<i>8.7</i>	<i>7.7</i>	<i>6.7</i>	19.0	18.0	17.0	16.0	15.0	14.0	13.0	12.0	11.0	10.0
23.0	<i>10.6</i>	<i>9.6</i>	<i>8.6</i>	<i>7.6</i>	19.0	18.0	17.0	16.0	15.0	14.0	13.0	12.0	11.0
24.0	<i>12.6</i>	<i>11.6</i>	<i>10.6</i>	<i>9.6</i>	<i>8.6</i>	19.0	18.0	17.0	16.0	15.0	14.0	13.0	12.0
25.0	<i>14.9</i>	<i>13.9</i>	<i>12.8</i>	<i>11.9</i>	<i>10.9</i>	<i>9.9</i>	19.0	18.0	17.0	16.0	15.0	14.0	13.0
26.0	<i>17.1</i>	<i>16.1</i>	<i>15.0</i>	<i>14.1</i>	<i>13.1</i>	<i>12.1</i>	<i>11.1</i>	19.0	18.0	17.0	16.0	15.0	14.0
27.0	<i>19.0</i>	<i>19.0</i>	<i>17.2</i>	<i>16.2</i>	<i>15.2</i>	<i>14.2</i>	<i>13.2</i>	<i>12.2</i>	19.0	18.0	17.0	16.0	15.0
28.0	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>18.5</i>	<i>17.5</i>	<i>16.5</i>	<i>15.5</i>	<i>14.5</i>	<i>13.5</i>	19.0	18.0	17.0	16.0
29.0	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>18.9</i>	<i>17.5</i>	<i>16.9</i>	<i>15.9</i>	<i>14.9</i>	19.0	18.0	17.0
30.0	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>18.0</i>	<i>17.0</i>	<i>16.0</i>	19.0	18.0
31.0	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>18.1</i>	<i>17.1</i>	19.0
32.0	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>
33.0	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>
34.0	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>
35.0	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>
36.0	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>
37.0	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>
38.0	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>	<i>19.0</i>

^aThese beliefs satisfy the monotonicity restriction (MR).

^bBeliefs off the equilibrium path are written in italics.

Table 6: Firm beliefs supporting pure strategy equilibria

e	c(e)	equilibrium effort level (e*)					
		0		1		2	
		$E[b[a y, e^*] e] - c(e)$	Diff.	$E[b[a y, e^*] e] - c(e)$	Diff.	$E[b[a y, e^*] e] - c(e)$	Diff.
0	0.00	10.00		9.00		8.11	
1	0.26	9.99	-0.01	9.74	0.74	8.79	0.69
2	0.57	9.98	-0.01	9.73	-0.01	9.43	0.64
3	0.93	9.98	-0.01	9.72	-0.01	9.42	-0.01
4	1.34	9.97	-0.01	9.71	-0.01	9.41	-0.01
5	1.80	9.96	-0.01	9.71	-0.01	9.41	-0.01
6	2.33	9.95	-0.01	9.70	-0.01	9.39	-0.01
7	2.92	9.94	-0.01	9.69	-0.01	9.38	-0.01
8	3.57	9.92	-0.02	9.72	0.03	9.37	-0.01
9	4.29	9.78	-0.14	9.64	-0.09	9.34	-0.04
10	5.08	9.52	-0.26	9.43	-0.21	9.18	-0.16
11	5.93	9.14	-0.38	9.10	-0.32	8.91	-0.27
12	6.84	8.65	-0.49	8.67	-0.44	8.52	-0.38
13	7.94	7.92	-0.73	7.99	-0.68	7.90	-0.63
14	9.23	6.95	-0.97	7.06	-0.92	7.03	-0.87
15	10.70	5.74	-1.21	5.91	-1.15	5.93	-1.10
16	12.35	4.30	-1.44	4.52	-1.39	4.59	-1.33
17	14.19	2.62	-1.68	2.89	-1.63	3.02	-1.58
18	16.22	0.70	-1.92	1.02	-1.87	1.20	-1.82
19	18.43	-1.46	-2.16	-1.08	-2.10	-0.86	-2.05

Table 7: Manager's effort decision for given equilibrium beliefs

e	c(e)	equilibrium effort level (e*)					
		3		4		5	
		$E[b[a y, e^*] e] - c(e)$	Diff.	$E[b[a y, e^*] e] - c(e)$	Diff.	$E[b[a y, e^*] e] - c(e)$	Diff.
0	0.00	7.32		6.63		6.05	
1	0.26	7.95	0.63	7.21	0.58	6.58	0.53
2	0.57	8.54	0.58	7.75	0.53	7.06	0.48
3	0.93	9.07	0.53	8.23	0.48	7.49	0.43
4	1.34	9.06	-0.01	8.66	0.43	7.87	0.38
5	1.80	9.05	-0.01	8.65	-0.01	8.2	0.33
6	2.33	9.04	-0.01	8.64	-0.01	8.19	-0.01
7	2.92	9.04	-0.01	8.64	-0.01	8.19	-0.01
8	3.57	9.03	-0.01	8.63	-0.01	8.18	-0.01
9	4.29	9.02	-0.01	8.62	-0.01	8.17	-0.01
10	5.08	8.91	-0.11	8.57	-0.05	8.16	-0.01
11	5.93	8.70	-0.22	8.40	-0.17	8.05	-0.11
12	6.84	8.37	-0.33	8.12	-0.28	7.82	-0.23
13	7.94	7.79	-0.57	7.60	-0.52	7.35	-0.47
14	9.23	6.98	-0.82	6.84	-0.76	6.64	-0.71
15	10.70	5.93	-1.05	5.84	-1.00	5.70	-0.94
16	12.35	4.64	-1.28	4.61	-1.23	4.52	-1.18
17	14.19	3.12	-1.52	3.14	-1.47	3.10	-1.42
18	16.22	1.35	-1.77	1.43	-1.71	1.44	-1.66
19	18.43	-0.65	-2.00	-0.52	-1.95	-0.45	-1.89

Table 8: Manager's effort decision for given equilibrium beliefs

e	c(e)	equilibrium effort level (e*)					
		6		7		8	
		$E[b[a y, e^*] e] - c(e)$	Diff.	$E[b[a y, e^*] e] - c(e)$	Diff.	$E[b[a y, e^*] e] - c(e)$	Diff.
0	0.00	5.58		5.21		4.92	
1	0.26	6.06	0.48	5.63	0.42	5.29	0.37
2	0.57	6.48	0.43	6.01	0.37	5.61	0.32
3	0.93	6.86	0.38	6.33	0.32	5.89	0.27
4	1.34	7.19	0.33	6.61	0.27	6.11	0.22
5	1.80	7.46	0.28	6.83	0.22	6.28	0.17
6	2.33	7.67	0.21	6.99	0.15	6.38	0.10
7	2.92	7.66	-0.01	7.08	0.09	6.42	0.04
8	3.57	7.66	-0.01	7.07	-0.01	6.43	0.01
9	4.29	7.65	-0.01	7.06	-0.01	6.42	-0.01
10	5.08	7.62	-0.03	7.06	-0.01	6.41	-0.01
11	5.93	7.56	-0.06	7.05	-0.01	6.41	-0.01
12	6.84	7.39	-0.17	6.93	-0.12	6.34	-0.07
13	7.94	6.97	-0.42	6.57	-0.36	6.03	-0.31
14	9.23	6.31	-0.66	5.96	-0.61	5.48	-0.55
15	10.70	5.42	-0.89	5.12	-0.84	4.69	-0.79
16	12.35	4.30	-1.12	4.05	-1.07	3.67	-1.02
17	14.19	2.93	-1.37	2.74	-1.31	2.41	-1.26
18	16.22	1.32	-1.61	1.18	-1.56	0.91	-1.50
19	18.43	-0.52	-1.84	-0.61	-1.79	-0.83	-1.74

Table 9: Manager's effort decision for given equilibrium beliefs

e	c(e)	equilibrium effort level (e*)					
		9		10		11	
		$E[b[a y, e^*] e] - c(e)$	Diff.	$E[b[a y, e^*] e] - c(e)$	Diff.	$E[b[a y, e^*] e] - c(e)$	Diff.
0	0.00	4.44		3.79		3.16	
1	0.26	4.76	0.32	4.06	0.27	3.37	0.21
2	0.57	5.02	0.27	4.28	0.22	3.54	0.16
3	0.93	5.24	0.22	4.44	0.17	3.65	0.11
4	1.34	5.41	0.17	4.56	0.12	3.74	0.09
5	1.80	5.53	0.12	4.68	0.12	3.83	0.09
6	2.33	5.62	0.09	4.75	0.08	3.91	0.08
7	2.92	5.67	0.05	4.82	0.07	3.96	0.06
8	3.57	5.70	0.03	4.88	0.06	4.01	0.05
9	4.29	5.71	0.01	4.91	0.03	4.05	0.03
10	5.08	5.70	-0.01	4.92	0.01	4.06	0.02
11	5.93	5.70	-0.01	4.91	-0.01	4.07	0.01
12	6.84	5.68	-0.02	4.90	-0.01	4.06	-0.01
13	7.94	5.42	-0.26	4.70	-0.21	3.91	-0.15
14	9.23	4.92	-0.50	4.25	-0.45	3.51	-0.40
15	10.70	4.19	-0.73	3.57	-0.68	2.88	-0.63
16	12.35	3.22	-0.97	2.66	-0.91	2.02	-0.86
17	14.19	2.02	-1.21	1.50	-1.16	0.92	-1.10
18	16.22	0.56	-1.45	0.10	-1.40	-0.43	-1.35
19	18.43	1.12	-1.68	-1.53	-1.63	-2.00	-1.58

Table 10: Manager's effort decision for given equilibrium beliefs

equilibrium effort level $e^* = 12$			
e	$c(e)$	$E[b[a y, e^*]] - c(e)$	Diff
0	0.00	1.47	
1	0.26	1.63	0.16
2	0.57	1.80	0.16
3	0.93	1.96	0.17
4	1.34	2.13	0.17
5	1.80	2.31	0.17
6	2.33	2.46	0.15
7	2.92	2.61	0.15
8	3.57	2.75	0.14
9	4.29	2.87	0.12
10	5.08	2.97	0.10
11	5.93	3.07	0.10
12	6.84	3.16	0.09
13	7.94	3.06	-0.10
14	9.23	2.72	-0.34
15	10.70	2.14	-0.58
16	12.35	1.33	-0.81
17	14.19	0.28	-1.05
18	16.22	-1.01	-1.29
19	18.43	-2.54	-1.53

Table 11: Manager’s effort decision with (MR) restriction on beliefs

C Instructions (translated from German)

Welcome to the experiment! In this experiment you can earn money. The amount of your earnings depends on your decisions and the decisions of the other participants. All decisions will be taken anonymously, i.e. nobody gets to know what decisions you have taken at any point. You will be paid at the end of the experiment. Payments are done one by one so that nobody gets to know the earnings of other participants. These instructions are the same for all participants. From now on, please refrain from talking to other participants and ask questions only to the experimenters directly.

Rounds and types of participants.

- The experiment lasts for twelve rounds.
- In the beginning of the experiment, you will be randomly assigned to one of two groups which consist of seven members each. You will belong to that group during the whole experiment. Neither during the experiment nor after will any participant get to know who else belonged to her group.
- In each group, there are two types of participants – A-participants and B-participants. There are three A-participants and four B-participants in each of the groups. In the beginning of the experiment, your type will be determined randomly and communicated to you. During the experiment, the roles of the participants remain unchanged.

The base value of A-participants. In the beginning of each round, each A-participant will be assigned a base value. This base value is drawn randomly from the numbers 0, 1, 2, 3, ... to 19 where each number is equally likely to occur. This base value remains the same throughout a round and will not be communicated to any participant (not even the A-participant herself). [RA-treatment: This base value remains the same throughout a round and will not be communicated to A-participants.] Note that each A-participant receives a new base value at the beginning of each new round.

Sequence of events during a round.

- A-participants
 - A base value from 0 to 19 is assigned to each A-participant. This base value cannot be observed by any participant.
 - Each A-participant chooses a number 0, 1, 2, 3, ... to 19. Each of these numbers is associated with a cost to the A-participant (see cost sheet).
 - The base value and the chosen number of the A-participant yield the so-called result for the A-participant. This result equals: base value + chosen number.
- B-participants
 - The results [RA-treatment: the results and the base values] of the three A-participants in a group are displayed to the B-participants in this group.
Note that the order of appearance of results on the screen is changed randomly in every round. It is therefore impossible to use the ranking of the results to infer the identity of an A-participant.
 - Each B-participant makes a transfer offer for each A-participant. This transfer offer has to lie between 0.00 and 38.00.
- A-participants
 - Each A-participant observes her chosen number, her result, and her transfer offers.
Note that the order of appearance of transfer offers on the screen is changed randomly in every round. It is therefore impossible to use the ranking of the transfer offers to infer the identity of a B-participant.
 - Each A-participant accepts one (or no) transfer offer.
- The payoffs per round are calculated and displayed.

Payoffs. The payoffs per round are calculated as follows.

- Payoff per round for A-participants.
 - If one transfer offer was accepted: payoff for this round = transfer offer – cost for chosen number.

- If no transfer offer was accepted: payoff for this round = $0 - \text{cost for chosen number}$.
- Payoff per round for B-participants.
 - If one transfer offer of the B-participant was accepted: payoff for this round = base value of the accepting A-participant – transfer offer to this A-participant
 - If several transfer offers of the B-participant were accepted: payoff for this round = sum of the base values of the accepting A-participants – sum of the transfer offers to these A-participants.
 - If no transfer offer of the B-participant was accepted, the payoff for this round is 0.

At the end of the experiment, the sum of the 12 round payoffs and an initial endowment will be exchanged into Euro (at a rate of 5 Euro-Cent per Taler) and paid out to you.

Please note: during the experiment, you will be asked to answer a few questions on the screen. Your answers to these questions will have no influence on your payments.