

PROJECT CHOICE AND SOCIAL IMAGE CONCERNS

Claudia Cerrone

Alessandro De Chiara

Ester Manna

Theo Saroglou

UB Economics *Working Paper No. 484***Title:** Project Choice and Social Image Concerns**Abstract:**

Employees' desire to impress their employer may lead to suboptimal choices, such as performing tasks that are out of their depth. In this paper, we formalise this intuition in a principal-agent setting and we experimentally analyse its practical relevance. Through a theoretical model, we show that an agent's desire to appear competent to their employer (social image concerns), can result in inefficient project selection. We test this prediction using a laboratory experiment and find that social image concerns increase the likelihood of suboptimal project choices when agents are male and the principal-agent interaction is not anonymous. Our findings have implications for organisational design.

JEL Codes: M5, C92, D82, D91**Keywords:** Social image, Delegation, Project choice, Gender.**Authors:**

Claudia Cerrone	Alessandro De Chiara	Ester Manna	Theo Saroglou
City St George's, University of London	Barcelona Economic Analysis Team (BEAT). Universitat de Barcelona.	Barcelona Economic Analysis Team (BEAT). Universitat de Barcelona.	Interdisciplinary Transformation University Austria
Email: Claudia.Cerrone@city.ac.uk	Email: aledeschiara@ub.edu	Email: estermann@ub.edu	Email: theo.saroglou@it-u.at

Date: April 2025

Acknowledgements: We would like to thank Olga Chiappinelli, Maria Cubel, Renaud Foucart, Georg Kirchsteiger, Christiane Schwierén, and Natalia Valdez Gonzalez for helpful comments. This work was financially supported by the Spanish Ministry of Science, Innovation and Universities, Spain (grant number PID2020-114040RB-I00), the Catalan Research Agency (grant number 2021SGR00678), and the European Union's Horizon Europe Research and Innovation Programme (grant Agreement number 101095175, SUSTAINWELL project). Views and opinions expressed are however those of the authors only and do not necessarily reflect those of the European Union.

Project Choice and Social Image Concerns*

Claudia Cerrone[†], Alessandro De Chiara[‡], Ester Manna[§] and
Theo Saroglou[¶]

April 10, 2025

Abstract

Employees' desire to impress their employer may lead to suboptimal choices, such as performing tasks that are out of their depth. In this paper, we formalise this intuition in a principal-agent setting and we experimentally analyse its practical relevance. Through a theoretical model, we show that an agent's desire to appear competent to their employer (social image concerns), can result in inefficient project selection. We test this prediction using a laboratory experiment and find that social image concerns increase the likelihood of suboptimal project choices when agents are male and the principal-agent interaction is not anonymous. Our findings have implications for organisational design.

Keywords: social image; delegation; project choice; gender.

JEL classifications: M5; C92; D82; D91.

*We would like to thank Olga Chiappinelli, Maria Cubel, Renaud Foucart, Georg Kirchsteiger, Christiane Schwieren, and Natalia Valdez Gonzalez for helpful comments. This work was financially supported by the Spanish Ministry of Science, Innovation and Universities, Spain (grant number PID2020-114040RB-I00), the Catalan Research Agency (grant number 2021SGR00678). and the European Union's Horizon Europe Research and Innovation Programme (grant Agreement number 101095175, SUSTAINWELL project). Views and opinions expressed are however those of the authors only and do not necessarily reflect those of the European Union.

[†]City St George's, University of London, Northampton Square, London EC1V 0HB, United Kingdom. E-mail: Claudia.Cerrone@city.ac.uk.

[‡]Department of Economics, Universitat de Barcelona and Barcelona Economic Analysis Team (BEAT), Avinguda Diagonal 696, 08034, Barcelona, Spain. E-mail: aledechiara@ub.edu.

[§]Department of Economics, Universitat de Barcelona and Barcelona Economic Analysis Team (BEAT), Avinguda Diagonal 696, 08034, Barcelona, Spain. E-mail: estermann@ub.edu.

[¶]Interdisciplinary Transformation University Austria, Altenberger Str. 66c/Science Park 4, 4040 Linz, Austria. E-mail: theo.saroglou@it-u.at.

1 Introduction

Evidence shows that people care not only about their material outcomes, but also about their social image, i.e., how they are perceived by others ([Andreoni and Bernheim, 2009](#), [Bursztyn and Jensen, 2017](#)). Social image concerns have been shown to affect a wide range of individual decisions, such as prosocial behaviours ([Lacetera and Macis, 2010](#), [Ariely et al., 2009](#)), welfare take-up ([Friedrichsen et al., 2018](#)), lying ([Bašić and Quercia, 2022](#)), and the intrinsic motivation to purchase ethical products ([Friedrichsen and Engelmann, 2018](#)). However, the effects of social image concerns on agents' behaviour in strategic settings have remained unexplored. We contribute to filling this gap by focusing on the role played by social image concerns in employer-employee interactions within organisations.

Due to digitalisation and advances in AI, employees are increasingly asked to carry out activities that fall outside their skill set ([Acemoglu et al., 2022](#)) and often accept tasks even if they do not have the appropriate expertise to perform them well.¹ This may happen because employees are unable to say no to a superior or because they want to try their luck given that successful completion of the task may generate material benefits (e.g., bonuses, promotions). It might also happen because of social image concerns: employees believe that by declining a task or requesting training, they would be perceived as unskilled by their employer (and colleagues). This behaviour has economically relevant implications. It can lead to poor performance, and thus bad outcomes for both employees and employers. Moreover, if employers anticipate this, they might be reluctant to delegate difficult or risky tasks to their employees.

In this paper, we theoretically and experimentally study whether social image concerns might induce employees to choose risky tasks even when they do not have the appropriate knowledge to perform them well. We also study whether this behaviour is foreseen by their employer and thus affects the employer's decision to delegate decision making authority to the employees. In our theoretical model, a principal (she) decides whether to delegate the choice of a risky project to a potentially better informed agent (he) or select a safe outside option. The probability that the chosen project is successful depends on the agent's ability. The agent privately observes his ability and chooses either a risky project or the safe outside option. If the agent has no or very weak social image concerns, there is no conflict of interest between the principal and the agent. This means that the agent will take the optimal action and the principal will always delegate decision making authority to the agent. However, if the agent has sufficiently strong social image concerns, a conflict

¹Even CEOs may accept tasks when they feel unprepared. In a survey of 402 CEOs from 11 countries, the leadership advisory firm Egon Zehnder found that more than two thirds of CEOs acknowledged that they were not fully prepared for the job and only half of them turned to senior management for advice ([Najipoor-Schutte and Patton, 2018](#)).

of interest may arise. The agent will suboptimally choose the risky project when he has low ability, which will adversely affect his own material payoff and that of the principal. This friction is exacerbated when the agent’s identity is more visible to the principal.

We test our model’s predictions through a laboratory experiment. Participants are randomly assigned to the role of principal or agent and keep their role throughout the experiment. Like in the model, principals must decide whether they want to delegate the choice of a risky project to their agent or select an outside option that gives both parties a sure payoff. Without knowing the principal’s delegation decision, agents must either choose one of three risky projects or select the outside option. Prior to making this choice, agents are informed about which project has the highest chance of succeeding only if they scored in the top 50% in a previously conducted IQ test. Payoffs are calibrated so that payoff-maximizing agents should choose a risky project if they know which project is best and should choose the outside option if they remain uninformed.

We have conducted three between-subject treatments. In the first treatment (*Baseline*), the principal does not find out whether an agent scored in the top 50% in the test. We use this as a baseline to determine whether some agents choose a risky project even if, due to their test performance, they did not learn which project was the best. To understand whether the decision to choose a project when uninformed is driven by social image concerns, we use two approaches. In our second treatment (*Principal learns*), we remove the agents’ possibility to signal their ability to the principal by choosing a risky project. We achieve this by letting the principal know whether the agent scored in the top 50% in the test. If the decision to choose a risky project despite being uninformed is driven by social image concerns, the fraction of uninformed agents choosing a risky project should be lower in this treatment than in the *Baseline*. In our third treatment (*Photo*), we take a more direct approach and make social image concerns stronger by removing the agents’ anonymity. This is achieved by showing the agent’s photo to the principal. As in the baseline, the principal does not find out whether the agent scored in the top 50% in the test. We expect the fraction of uninformed agents choosing a risky project to be higher in this treatment than in the *Baseline*. The experiment is deliberately designed so as to shut down motives different from social image concerns, like the inability to say no to a superior, or the desire to qualify for promotions or bonuses.

We find that in the *Baseline* nearly half of the agents (48%) choose a risky project even when they do not know which project is best, i.e., even when their expected monetary payoff would have been higher had they selected the outside option. This fraction remains nearly identical in the *Principal learns* treatment (47%), but increases to 64% in the *Photo* treatment. Notably, the latter increase is driven by male agents. The fraction of uninformed male agents who choose a project increases from 32% and 40% in the two anonymous treatments (*Baseline* and *Principal learns* treatment, respectively) to 71% in

the non-anonymous treatment (*Photo* treatment). In contrast, the fraction of uninformed female agents who choose a risky project does not significantly vary across treatments. As real-life workplace interactions are typically not anonymous, our experimental findings suggest that social image concerns can cause agents to make suboptimal decisions to signal their ability to the principal. Our results also highlight strong gender differences in social image concerns.

The rest of the paper proceeds as follows. In Section 2, we review the related literature. In Section 3, we develop the theoretical framework. In Section 4, we describe the experimental design. In Section 5, we present the results of our experiments. In Section 6 we further discuss our findings and in Section 7 we provide concluding remarks.

2 Related Literature

Our paper relates to the literature on the impact of social image concerns on behaviour. [Friedrichsen and Engelmann \(2018\)](#) use an experiment to investigate the link between social approval and the intrinsic motivation to purchase ethically. They find that participants are more willing to pay a premium for Fairtrade chocolate when their choice can be observed by others. This effect is driven by participants who had not chosen Fairtrade chocolate in a pre-lab choice, highlighting the influence of social image concerns on ethical behaviour. [Friedrichsen et al. \(2018\)](#) use an experiment to study the impact of social image concerns on the take-up of welfare benefits. They find that participants avoid claiming benefits to evade the stigma associated with being perceived as low-skilled or dependent on others. This highlights that social image concerns can lead to economically suboptimal decisions, as individuals prioritise maintaining a positive social image over receiving monetary transfers. [Lacetera and Macis \(2010\)](#) use the entire population of blood donors in an Italian town to study the impact of social image concerns on prosocial behaviour. They find that donors increase their donations immediately before reaching the thresholds for which rewards are given, but only if the prizes are publicly announced. This shows that social image concerns are an important driver of prosocial behaviour. [Bašić and Quercia \(2022\)](#) study the influence of self and social image concerns on lying behaviour. They find evidence supporting social image concerns as a source of lying costs, but no evidence supporting self image concerns. While the aforementioned papers focus on the impact of social image concerns on *individual* decisions, our paper is the first to study the role played by social image concerns in a *strategic* environment - the interaction between a principal and an agent.

Our paper is also related to [Katok and Siemsen \(2011\)](#), who conduct an experiment to study the impact of reputation concerns on task choice. They find that when reputation

concerns are present, both highly capable and less capable agents choose more difficult tasks, leading to reduced success rates. When reputation concerns are removed, success rates become significantly higher. However, the treatment that eliminated reputation concerns also eliminated social image concerns, as well as any incentives for highly capable and less capable agents to choose a hard task. As a result, the increase in success rates might have been driven not only by the absence of reputation concerns but also by the elimination of social image concerns.

Our findings contribute to the literature on gender differences in social image concerns. [Murad et al. \(2019\)](#) explore the impact of competitive, social-value, and social image incentives on task performance and examine gender differences in response to these incentives. They find that competitive incentives and social image incentives significantly improve performance, with competitive incentives having a more pronounced effect on males. No significant gender difference was observed in the social image incentives group.

Methodologically, our study relates to the literature that studies the effects of removing anonymity on individual and group behaviour. In the field, [Bursztyn and Jensen \(2015\)](#) show that non-honor students in disadvantaged areas are significantly less likely to sign up for free access to a SAT preparation course when this decision would be revealed to their classmates. The interpretation provided is that these students dislike signaling to their classmates that they care about going to college, since this could be perceived as a departure from a social norm. In the lab, [Guo and Recalde \(2023\)](#) find that voicing disagreement occurs less frequently when the interaction is not anonymous. In their theoretical framework they rationalise this result by positing that social image concerns exacerbate the disutility of voicing disagreement. They also find that participants are more likely to override women’s rather than men’s early choices. In our experiment, we document how male agents are more likely to make suboptimal choices when anonymity is removed so that social image concerns are magnified.

3 Theoretical Model

A principal (she) must decide whether she delegates the choice of a risky project to an agent (he) or selects a safe outside option. The probability that the project is successful depends on the agent’s ability. The agent privately observes his ability and either chooses a project or selects the outside option. Real-world applications abound: from investors who may ask their financial advisors about investment strategies to CEOs who may delegate the choice of revamping a product line to their employees. We extend this standard delegation and project choice setting by allowing for social image concerns. Specifically, we assume that the agent may care about being viewed as of high ability by the principal.

Both the agent and the principal are assumed to be risk neutral and self-interested.

We let θ denote the agent's type, where $\theta \in \{L, H\}$, and $\gamma \in (0, 1)$ denotes the probability that the agent is of type H . The agent knows his own type and chooses an action $x \in \{a, \emptyset\}$, which affects an outcome $y \in \{S, F, R\}$. Action a is risky, as it can generate either $y = S$ or $y = F$. It holds that

$$Pr[y = S | \theta = H] = p > q = Pr[y = S | \theta = L],$$

with $p < 1$ and $q > 0$. Action $\emptyset$ is safe, as it generates a type-independent outcome $y = R$ for sure.

Let $m(x, \theta)$ be the expected monetary payoff to a type- θ agent who chooses action x . It holds that

$$m(a, H) > m(\emptyset, H) = m(\emptyset, L) > m(a, L).$$

In particular, we assume that if $y = S$ (respectively, $y = R$) the principal receives s_P (r_P) and the agent receives s_A (r_A). If $y = F$ both the principal and the agent receive a payoff normalised to 0. As a result, the above assumption on the agent's expected monetary payoffs implies that

$$ps_A > r_A > qs_A > 0. \quad (1)$$

For the principal we assume the following inequalities:

$$ps_P > r_P > qs_P > 0. \quad (2)$$

The principal does not know the agent's type, but she observes which action the agent has taken and which outcome has realised. She can form a belief about the agent being a high type, which we denote as $\tau(y) \in [0, 1]$, as it depends on the outcome $y \in \{S, F, R\}$. The principal's belief about the agent's type affects his *reputational payoff*. The agent's sensitivity to social image concerns is measured by the parameter η . The expected utility of agent θ with sensitivity η can be written as:

$$U_{\theta, \eta}(x) = m(x(\theta, \eta), \theta) + \eta \cdot [\tau(y(x(\theta, \eta), \theta))H(\gamma) + [1 - \tau(y(x(\theta, \eta), \theta))]L(\gamma)],$$

where $H > 0$ is the sense of *pride* the agent feels when he is perceived as a high type and $L < 0$ is the *shame* that he feels when he is perceived as a low type. Plausibly, these feelings of pride and shame depend on the proportion of high versus low-ability individuals in the population. Accordingly, we posit that $\frac{\partial H}{\partial \gamma} \leq 0$ and $\frac{\partial L}{\partial \gamma} \geq 0$. In words, the feeling of pride (respectively, shame) is non-increasing (non-decreasing) in the fraction of high-type agents in the population. The agent's reputational payoff is increasing in the likelihood that the principal thinks he is a high type $\tau(y)$. We assume that η is known to be distributed according to a continuous distribution function $G(\cdot)$ with density $g(\cdot)$ on $[0, \infty)$.

Timing of the model. The sequence of events is as follows.

1. Nature draws the agent's type.
2. The principal decides whether to delegate the choice of the project to the agent ($d = 1$) or select the outside option ($d = 0$).
3. The agent observes his type and, if $d = 1$, decides which action x to take.

3.1 Equilibrium Analysis

We look for the Bayesian Equilibria of the game, which are characterised by (i) the agent's project choice, x , as a function of his ability, θ , and social-image concerns, η ; (ii) the principal's belief about the agent's type, which is a function of the outcome, $\tau(y)$; (iii) the principal's delegation choice, d , which is a function of the expected agent's ability and social image concerns. Being rational, the agent will correctly anticipate the principal's belief about his type, which is formed according to Bayes' rule.

We begin by considering the scenario with no social-image concerns.

Remark 1. *If $\eta = 0$, the agent's action is $x(H) = a$, $x(L) = \emptyset$, and the principal's belief is $\tau(S) = \tau(F) = 1$ and $\tau(R) = 0$. The principal always delegates, i.e., $d = 1$.*

If the agent does not attach any value to his reputational payoff, he will choose the risky action when he is a high type and the outside option otherwise. Therefore, choosing a project would always reveal that the agent is a high type, irrespective of the outcome, and choosing the outside option would always reveal that the agent is a low type. Thus, in the absence of social image concerns, there is no conflict of interest between the principal and the agent. The principal always delegates because she can trust the agent will always take the optimal action.

We now study the more general case where η can be positive. To make the problem interesting we assume that if all L -type agents choose $x = a$, the principal would rather select the outside option than delegate this choice to the agent:

$$\gamma p_{SP} + (1 - \gamma)q_{SP} < r_P. \quad (3)$$

The first result we prove is that high-type agents never want to opt for the outside option: the risky option yields a higher monetary reward in expectation and allows them to signal higher ability.

Lemma 1. *H-type agents always choose $x = a$.*

Social image concerns do not alter the behaviour of H-type agents who continue to select the risky task. However, there cannot be a fully-separating equilibrium because a low-type agent who cares enough about his social image would rather deviate to be viewed as a high type.

Lemma 2. *There cannot be a fully-separating equilibrium where all agents with $\theta = H$ choose $x = a$ and all agents with $\theta = L$ choose $x = \emptyset$.*

In the next proposition, we characterise the cutoff value of η above which an L-type agent chooses the risky action.

Proposition 1. *There exists a threshold value $\hat{\eta} > 0$, such that all L-type agents with $\eta > \hat{\eta}$ choose $x = a$ and all L-type agents with $\eta < \hat{\eta}$ choose $x = \emptyset$.*

L-type agents with strong enough social image concerns will be willing to suboptimally choose a risky action.

We now proceed to investigate the principal's choice. In equilibrium, the principal correctly predicts the behaviour of the different types of agents. Hence, she is more likely to delegate authority over the choice of the project to the agent when the proportion of H-types relative to L-types (γ) is higher, when the proportion of low types who choose the safe option ($\hat{\eta}$) is higher, and when her personal gain from a successful implementation of the task ($s_P - r_P$) is greater. The following proposition formalises the principal's delegation decision under the assumption that, if indifferent, she prefers to delegate authority.

Proposition 2. *The principal chooses $d = 1$ if $\hat{\eta} \geq \bar{\eta}$ and $d = 0$ otherwise, where $\bar{\eta}$ is determined from the following equation:*

$$G(\bar{\eta}) = 1 - \left(\frac{\gamma}{1 - \gamma} \right) \left(\frac{ps_P - r}{r - qs_P} \right). \quad (4)$$

Plausibly, an agent's concern for social image will be stronger when the agent's identity is more visible to his principal. For example, within organisations employees are usually well known to their superiors and regularly interact with them. Thus, visibility is high. However, in other principal-agent settings (e.g., when agents are hired as freelancers via online platforms), the agent's identity is less visible. To take this into account, we modify the agent's utility by assuming that the agent's reputational payoff depends on a parameter $V \geq 0$, which stands for visibility. That is, an agent's expected utility is:

$$U_{\theta, \eta}(x) = m(x(\theta, \eta), \theta) + \eta \cdot \left[\tau(y(x(\theta, \eta), \theta))H + [1 - \tau(y(x(\theta, \eta), \theta))]L \right] V.$$

The next proposition studies the effect of visibility on the agents' and the principal's choice.

Proposition 3. *An increase in visibility V induces more L -type agents to choose $x = a$ and, accordingly, increases the likelihood that a principal chooses $d = 0$*

When the agent is more visible to the principal, a risky action is more likely to be undertaken by the low-ability agent and the principal will be less inclined to delegate. When the agent is less known or not known to the principal, the principal-agent friction may disappear, even if the agent is generally concerned about his social image. In the appendix, we provide additional results concerning the effect of different parameter values on the choice variables.

Related models. Our model of delegation of authority builds on [Aghion and Tirole \(1997\)](#). There are some noteworthy differences, though. In our model the principal cannot implement the task, and the agent’s knowledge of the projects’ payoffs depends on his ability and not on his effort choice. Moreover, in their model the principal and the agent have material interests in undertaking different projects, whereas in ours social image concerns are the source of the conflict of interests. Akin to many standard delegation models, we also assume that the principal cannot commit to contingent transfers (e.g., see [Alonso and Matouschek, 2008](#)).

In modelling social image concerns, we follow [Bénabou and Tirole \(2006\)](#) and [Ellingsen and Johannesson \(2008\)](#). Our model also relates to [Khalmetski and Sliwka \(2019\)](#) where agents do not want to be viewed as liars. However, unlike us, they study individual decision makers who incur a lying cost and their choice set is not necessarily binary.

4 Experimental Design

The experiment uses a between-subject design and has three parts.

Part 1. In the first part of the experiment, each participant is randomly assigned to the role of agent or to the role of principal, neutrally termed as “role A” and “role B”, respectively. Participants are randomly matched in groups of five members: one principal and four agents.² This was illustrated to participants by using Figure 1. Every agent is given up to 10 minutes to complete a 30-question IQ test. Principals do not take the test but can view the questions if they want to.

Part 2. In the second and core part of the experiment, every participant is told that there are three boxes: one blue, one green, and one red. Two boxes are empty. When an

²We match each principal with multiple agents as we are primarily interested in the behaviour of the agents.

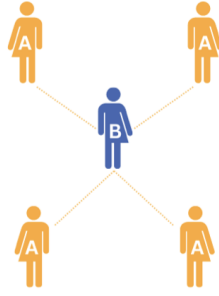


Figure 1: Group

empty box is chosen, both the principal and the agent earn €0. The other box contains some money and, when chosen, with a 75% probability it can be opened and will pay €12 to the agent and €10 to the principal, and with a 25% probability it cannot be opened and will leave both the agent and the principal with €0. If no box is chosen (outside option), the principal earns €6 and the agent earns €4. The box that contains the money may be different across agents.

Initially, no participant knows which of the three boxes contains the money. They only know that each box has the same probability of containing the money. Later in the experiment, depending on their relative performance in the IQ test, the agents might find out which box contains the money, as further explained below. Principals do not find out which box contains the money.

After receiving the instructions and answering some review questions, participants go through the three following stages.

Stage 1: Each principal must decide between choosing the outside option and allowing all the agents she is linked to (henceforth referred to as “her agents”) to make a decision. If a principal delegates the decision, one of her four agents will be selected at random and their decision will determine the principal’s payoff.

Stage 2: Each agent learns whether they scored in the top 50% in the test among the participants who took the test in the session. If they did, they immediately find out which box contains the money. If they did not, they do not find out. Potential ties are broken by using the agents’ completion time in the test.

Stage 3: Without knowing the decision made by their principal in Stage 1, each agent must choose either the outside option or one of the boxes. Their principal will observe this decision and the resulting payment. Depending on the treatment (see below), the principal might also find out whether the agent scored in the top 50% in the test. Note that an agent’s decision will only be implemented if their principal delegated decision making authority to them in Stage 1. If a principal chose the outside option, her agents’ decisions will have no consequences.

Part 3. At the end of the experiment, we ask agents whether they want to know their percentage of correct answers in the test and their ranking relative to other participants. We also ask all participants some final questions to elicit their attitudes towards risk (Charness and Gneezy, 2010), their social preferences (Bartling et al., 2009), and their demographics. Finally, we ask participants how many other participants they knew in their session and, when relevant, we ask the principal how many agents they knew in their group. The experiment’s instructions can be found in Appendix C.

Treatments. Our theoretical model predicts that agents with sufficiently strong social image concerns will choose a risky project even if they do not know which project is best (Proposition 1). The fraction of uninformed agents choosing a risky project will be higher when the agent’s visibility is higher (Proposition 3). The key goal of our experiment is twofold. First, we want to determine whether uninformed agents are more likely to choose a risky project when this choice can signal high ability to the principal as opposed to when this choice cannot be used to signal high ability (as the principal learns about the agent’s true ability). Second, we want to determine whether an increase in visibility prompts a higher fraction of uninformed agents to choose a risky project. To achieve these goals we use three treatments. In our first treatment (*Baseline*), the principal does not find out whether her agents scored in the top 50% in the test. We use this as a baseline to establish whether a fraction of agents chooses a box even if uninformed. In our second treatment (*Principal learns*), we inform the principal about her agents’ relative performance in the test. This effectively removes the possibility that uninformed agents choose a project to signal their ability to the principal. That is, social image concerns can no longer affect the agents’ decisions. Visibility in these two treatments is very low, as the principal-agent interaction is fully anonymous. However, agents might still care about how they are viewed by the principal. Our hypothesis is the following. If uninformed agents choose a box because of social image concerns, then the fraction of uninformed agents choosing a box should be lower when the principal finds out about her agents’ relative performance than in the baseline where the principal does not.

Hypothesis 1. *The proportion of uninformed agents selecting a risky project will be significantly lower when the principal learns the agents’ relative test performance than when she does not.*

In our third treatment (*Photo*), the principal does not find out about her agents’ relative performance in the test as in the *Baseline*, and we remove the agents’ anonymity by showing each agent’s photo to their principal. Specifically, the agents’ pictures are shown to their principal together with information about the choices they made and the

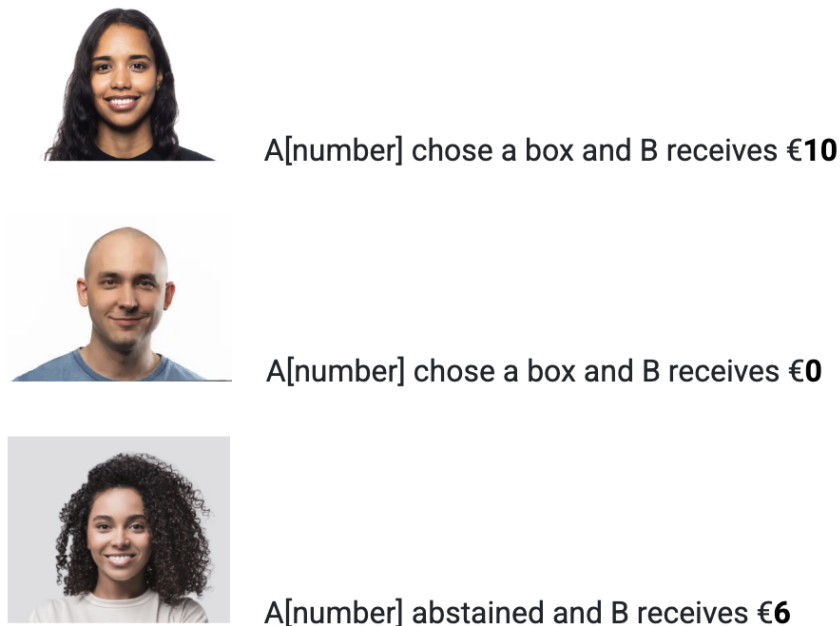


Figure 2: What a principal observes in the *Photo* treatment (illustration).

corresponding payoffs.³ Figure 2 provides an illustration of what the principal observed, as shown to participants in the experimental instructions. Recall that, according to Proposition 3, an increase in visibility induces more low ability agents to choose a risky project. Based on Proposition 3, our hypothesis is that the removal of anonymity will increase agents' social image concerns and thus the fraction of uninformed agents choosing a box.

Hypothesis 2. *The proportion of uninformed agents choosing a risky project will be significantly higher when the agents' photos are shown to their principal than when the principal-agent interaction is fully anonymous.*

It is worth remarking that the agents do not observe their principal's picture. We deliberately opted for this to avoid potential confounding factors arising from the agents observing the principal's picture, such as gender or preexisting friendships, over which we had limited control.

Table 1 summarises our treatments.

Implementation. The experimental sessions were programmed in o-Tree and run at the lab of the University of Barcelona in 2024. In total, 300 subjects participated in the experiment: 240 were assigned to the role of agent ("role A") and 60 to the role of principal ("role B"). The average duration of the experimental sessions was about 35 minutes. The

³In this treatment, all participants had to take a picture using the computer webcam right after the end of Part 1, i.e., after the completion of the test but before reading the instructions for Part 2.

Table 1: Treatments

Treatment	Principal learns agents' performance	Agents' identity is anonymous	# subjects
Baseline	No	Yes	100
Principal learns	Yes	Yes	75
Photo	No	No	125

average earnings were €11, including a show-up payment of €5. In each session, there were between one and four groups (i.e., between 5 and 20 participants). Prior to starting the experimental sessions, the experiment was preregistered on AsPredicted (preregistration number 170427).

5 Experimental Results

In this section, we report the results of our experiment. Subsection 5.1 focuses on the agent's decision - the main focus of the paper. Subsection 5.2 briefly discusses the principal's decision.

5.1 The agent's decision

Our main goal is to determine whether social image concerns increase the fraction of uninformed agents who choose a risky project. To this aim, we vary across treatments (i) whether the principal learns about their agents' relative performance in the test, and (ii) whether the agent's identity remains anonymous to the principal.

Table 2: Frequency of agents choosing a box by treatment

Treatment	Agent is uninformed	Agent is informed
Baseline	.48 (40)	1 (40)
Principal learns	.47 (30)	.93 (30)
Photo	.64 (50)	.9 (50)

Number of observations in parentheses.

Frequency of project choice. Table 2 illustrates the agents' frequency of choosing a project (in the experiment, a box) over the outside option by treatment. In the *Baseline*,

where the principal does not find out the agents’ relative performance in the test, nearly half of the agents (48%) choose a box even if uninformed. All the informed agents choose a box, which reassures us that participants understood the task and responded to incentives properly.

In the *Principal learns* treatment, where uninformed agents can no longer signal their ability to their principal by choosing a box, we still observe that nearly half of them (47%) choose a box. The fraction of uninformed agents choosing a box is nearly identical, regardless of whether their principal learns about their relative performance in the test or not. Hence, we find no evidence supporting Hypothesis 1.

A potential reason why social image concerns do not seem to play a role in these two treatments is that the agents’ interaction with their principal remains fully anonymous. In contrast, in the *Photo* treatment anonymity is lifted: each agent’s photo is shown to their principal. The fraction of uninformed agents choosing a box rises to 64%, which is significantly higher than the fraction of uninformed agents choosing a box in the two non-anonymous treatments (one-tailed t-test, $p = 0.034$). This supports Hypothesis 2.

Table 3: Frequency of uninformed agents choosing a box by treatment and gender

Treatment	Male	Female	Total
Baseline	.32 (22)	.67 (18)	.48 (40)
Principal learns	.4 (20)	.6 (10)	.47 (30)
Photo	.71 (24)	.58 (26)	.64 (50)

Number of observations in parentheses.

Next, we examine whether uninformed agents’ behaviour varies by gender. Table 3 illustrates uninformed agents’ frequency of choosing the box by both treatment and gender.⁴ We observe strong gender differences. The fraction of uninformed male agents who choose a box in the *Photo* treatment is significantly higher than the fraction of uninformed male agents who choose a box in the two anonymous treatments (one-tailed t-test, $p = 0.0028$). In contrast, there is no significant difference for female agents.⁵ In Figure 5.1, we provide a graphical representation of these results. The left panel displays the fraction of all uninformed agents who chose a box in the two anonymous treatments

⁴In the male category, there is also one uninformed agent who classified themselves as “Other”. Our results do not change if we include this observation in the female category or if we drop it.

⁵We also observe that, while more uninformed males than females choose a box in the *Photo* treatment, the opposite is true in the *Baseline* and in the *Principal learns*.

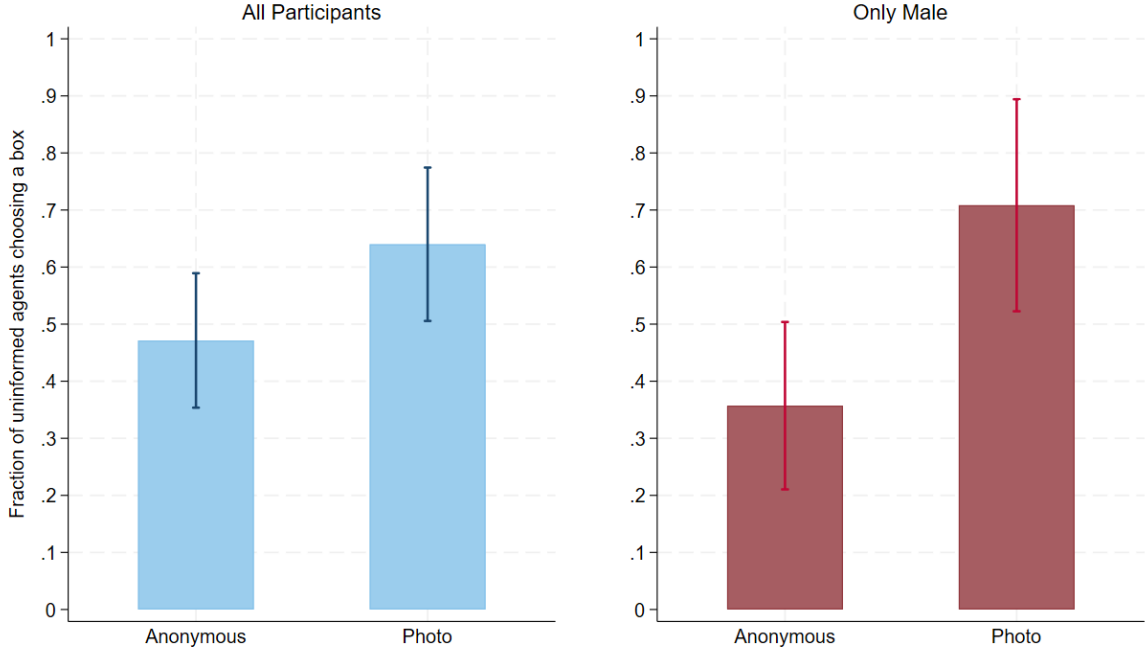


Figure 3: Fraction of uninformed agents choosing a box.

(47%) compared to the *Photo* treatment (64%). The right panel only shows the fraction of *male* uninformed agents choosing a box, across the same treatment conditions (37% versus 71%). We defer a discussion of these results and their managerial implications to the next sections.

The effect of the principal learning about the agents' performance on project choice. To explore how the uninformed agents' project choice is affected by their principal finding out about the agents' relative test performance, we perform a logit estimation of the probability that uninformed agents choose a box. The dependent variable, *Agent chooses box*, takes value 1 if an uninformed agent chooses a box and 0 if he selects the outside option. The explanatory variables we consider are the following. The treatment variable *Principal learns* takes value 1 if the principal learns whether each of their agents scored in the top 50% in the test and 0 if the principal does not learn. The variable *Female* takes value 1 if the agent is female and 0 otherwise. Our basic specification, Specification (1), controls for *Principal learns*, *Female* and their interaction.

We also have other three specifications where we include additional explanatory variables. The variable *Investment* measures subjects' attitude towards risk (the higher the variable, the lower the subject's risk aversion).⁶ *Prosocial* is a dummy that takes value 1

⁶This variable measures the fraction of a subject's endowment that they decide to invest, where the amount invested is equally likely to be lost or to become 2.5 times larger.

if a participant is both weakly and strongly prosocial according to the post-experimental test to elicit social preferences.⁷ The variable *People known* denotes the number of people present in the room that a participant reports they know. Finally, the variable *See score* takes value 1 if an agent decides to learn their score and ranking in the test at the end of the experiment and 0 otherwise. Specification (2) is the same as Specification (1) but also controls for *Investment*. Specification (3) additionally controls for *Prosocial*, and Specification (4) additionally controls for *People known* and *See score*.

The results of this logit regression can be found in Table 7 in Appendix B.2. Table 4 reports the marginal effects of informing the principal about their agents' test performance on uninformed agents' probability of choosing a project. Given the gender effects shown in Table 3, we report separate marginal effects for female and male agents. As in the tests, we do not observe a treatment effect. Informing the principal about their agents' test performance does not affect uninformed agents' project choice.

Table 4: Marginal effects of informing the principal on uninformed agents' choice

DV: Agent chooses box				
	(1)	(2)	(3)	(4)
Marginal effect of <i>Principal learns</i>:				
if Female=0	0.082	0.090	0.073	0.039
	(0.149)	(0.150)	(0.147)	(0.155)
if Female=1	-0.067	-0.081	-0.056	-0.113
	(0.192)	(0.190)	(0.195)	(0.191)
N_I	70	70	70	70

*** $p < 0.01$; ** $p < 0.05$; * $p < 0.1$

Logit regression. Standard errors in parentheses. All coefficients are average marginal effects of informing the principal about their agents' test performance on uninformed agents' probability of choosing a box. N_I denotes the number of uninformed agents in the *Baseline* and the *Principal learns* treatment.

The effect of lifting anonymity on uninformed agents' project choice To analyse the impact of showing the agents' photos to their principals on uninformed agents' project choice, we use a logit regression. The dependent variable is *Agent chooses box*. The explanatory variables we consider are the same as in Table 7 in Appendix B.2, with one difference. The treatment variable, *Not anonymous*, takes value 1 if an agent is in the

⁷A participant is weakly prosocial if, when faced with two allocations where they earn the same amount, they prefer the allocation where their partner earns more. A participant is strongly prosocial if they prefer to equally split a pie with their partner than to have a bigger share of the pie than the other.

Photo treatment (i.e., if the agents' photos are shown to the principal) and 0 if he is in the *Baseline* or *Principal learns* treatment (i.e., if agents' identity remains fully anonymous).

Specification (1), controls for *Not anonymous*, *Female* and their interaction. Specification (2) is the same as Specification (1) but also controls for *Investment*. Specification (3) additionally controls for *Prosocial*, and Specification (4) additionally controls for *People known* and *See score*. Note that, while there is a difference between the two anonymous treatments (in one the principal learns about their agents' performance and in the other she does not), it is reasonable to pool their data into a joint *Anonymous* treatment as we observed no difference in behaviour between the two.⁸

The results of the logit regression can be found in Table 8 in Appendix B.2. Table 5 reports the marginal effects of showing each agent's photo to their principal on uninformed agents' probability of choosing a project. We find that when agents' photos are shown to their principal, the fraction of uninformed male agents choosing a box significantly increases. For female agents, there is no significant effect. Our results are robust to different specifications.⁹

Table 5: Marginal effect of lifting anonymity on uninformed agents' project choice

DV: Agent chooses box				
	(1)	(2)	(3)	(4)
Marginal effect of <i>Not anonymous</i>:				
if Female=0	0.351*** (0.119)	0.372*** (0.115)	0.371*** (0.117)	0.376*** (0.119)
if Female=1	-0.066 (0.133)	-0.064 (0.125)	-0.070 (0.120)	-0.039 (0.119)
N_I	120	120	120	120

*** $p < 0.01$; ** $p < 0.05$; * $p < 0.1$

Logit regression. Standard errors in parentheses. All coefficients are average marginal effects of showing each agent's photo to their principal on uninformed agents' probability of choosing a box.

N_I denotes the number of uninformed agents in our three treatments.

⁸Recall that, as shown in Table 2, the fraction of uninformed agents choosing a box is nearly identical in the two anonymous treatments.

⁹In Table 9 in Appendix B, we report the aggregate marginal effects of showing the agents' photos to their principal on uninformed agents' probability of choosing a project, without distinguishing between males and females. The effect is still statically significant.

5.2 The principal’s decision

While this paper’s focus is on the agent’s project choice, it is interesting to look at the principal’s delegation decision too. Table 6 reports the frequency of delegation by treatment and gender. In aggregate, principals delegate significantly more in the *Photo* treatment than in the two anonymous treatments (*Baseline* and in the *Principal learns*). This suggests that principals do not anticipate that the *Photo* treatment will increase social image concerns and thus make the choice of the box more likely. A potential explanation is that when principals are shown their agents’ identity, they feel more trustful. When we split the principals by gender, we find that this effect is driven by females. Female principals delegate significantly more in the *Photo* treatment than in the two anonymous treatments (one-tailed t-test, $p = 0.0011$). For male principals there is no significant difference.

Female principals delegate more than male principals in the *Photo* treatment (one-tailed t-test, $p = 0.0160$), and less than male principals in the two anonymous treatments, but the latter difference is only marginally significant (one-tailed t-test, $p = 0.0576$).

Table 6: Frequency of delegation by treatment and gender

Treatment	Male	Female	Total
Baseline	.818 (11)	.556 (9)	.7 (20)
Principal learns	.857 (7)	.625 (8)	.733 (15)
Photo	.75 (7)	1 (17)	.92 (25)

Number of observations in parentheses.

6 Discussion

In this section, we further discuss the results described in Section 5.

Why uninformed agents choose a project in the anonymous treatments. As discussed in Section 5, the fraction of uninformed agents who choose a project is nearly identical in the *Baseline* and in the *Principal learns* treatment, even if only in the latter the principal is informed about the agents’ performance in the test. This suggests that the agents’ decision to choose a box even if uninformed, observed in the *Baseline*, is not driven by their desire to signal their ability to the principal. What is driving this behaviour then? In what follows, we discuss some potential drivers.

A first potential driver are risk attitudes. Agents might choose a project (risky action) even if uninformed because they are risk lovers. A second potential justification is that agents overweigh small probabilities of winning. A third potential reason is that agents might derive “ego utility” from maintaining a positive self-image. [Kőszegi \(2006\)](#) shows that individuals might prefer more challenging tasks to enhance their self-image, even when these tasks do not match their actual skills and/or are not materially beneficial for them. This can lead to overconfidence and misaligned task choices. In a similar vein, in our experiment, agents might decide to pick a project (the more interesting and active decision) rather than the outside option (passive decision) as it will make them feel better about themselves. A fourth potential explanation is that agents are inequity averse and dislike disadvantageous inequity more than they dislike advantageous inequity. By selecting the outside option agents earn less than the principal, while if they choose a project and the project is successful they will earn more than the principal. Hence, agents might choose a project even if uninformed as they do not want to earn less than the principal. To test for this, we use one of the post-experimental questions measuring social preferences ([Bartling et al., 2009](#)) to create a variable that measures agents’ dislike for disadvantageous inequality. We find that this variable has no significant effect on uninformed agents’ probability to choose a risky project. It should be noted that all the potential factors aforementioned, if present, are orthogonal to our main treatment variation (anonymous interaction versus *Photo* treatment) and thus do not affect our main result discussed in Section 5.

Gender effects in social image concerns. Our result that men exhibit greater social image concerns than women aligns with established research showing that men are more responsive to competition ([Gneezy et al., 2003](#), [Gneezy and Rustichini, 2004](#), [Backus et al., 2023](#)). This finding is also in line with existing evidence on overconfidence and social image concerns. In the experiment by [Ewers and Zimmermann \(2015\)](#), participants answer quiz questions and report their performance either privately or publicly. Public reports are significantly inflated compared to private ones, suggesting that social image concerns lead to overconfident self-assessments. Notably, men exhibit a stronger tendency to inflate their performance in the presence of an audience than women, indicating stronger social image concerns.

Money left on the table. In the previous section, we have documented that agents make suboptimal decisions, especially in non-anonymous interactions. It is valuable to measure the derived payoff loss. To establish a benchmark, we calculate the potential gain from delegation as the difference between the maximum achievable joint payoff through delegation and the joint payoff when the principals decide not to delegate. We find that

the realised payoffs represent 60%-67% of the potential gain from delegation in the two anonymous treatments, and only 46% of the potential gain in the *Photo* treatment. This shows that when anonymity is lifted and uninformed agents become more likely to choose a project, the implied payoff loss gets substantially larger. The computations are reported in Appendix [B.1](#).

Willful ignorance. When agents do not score in the top 50% in the test, they are significantly less likely to decide to know their score and ranking at the end of the experiment (one-tailed t-test $p=0.023$). This effect is more pronounced among male agents (one-tailed t-test $p=0.012$), suggesting the presence of self-image concerns.¹⁰ This observation is consistent with recent experimental findings indicating that individuals frequently avoid performance feedback, particularly when they anticipate negative outcomes ([Petrishcheva, 2023](#)).

7 Conclusions

While social image concerns have been shown to affect a wide range of individual decisions, its effects on strategic behaviour have remained unexplored. In this paper, we investigate the role played by social image concerns in employer-employee interactions within organisations. Due to upskilling, employees are increasingly asked to perform tasks for which they might not yet have adequate skills. Some employees might accept tasks even if they lack the necessary knowledge to perform them well. One potential reason for this is social image concerns: employees may hesitate to decline a responsibility or request training or support, as doing so could signal low ability to their superior.

Our paper offers both theoretical and empirical contributions. Using a theoretical model, we show that social image concerns can create a conflict of interest between a principal and an agent, leading the agent to perform a task even when he does not have the necessary knowledge to perform it well. This effect is stronger when the task is more visible. Using a laboratory experiment, we find that social image concerns affect agents' behaviour when the interaction is not anonymous. This is in line with our model's predictions. Notably, the effect is fully driven by male participants, who are almost twice as likely to choose a project when uninformed if the interaction is not anonymous as compared to when it is anonymous.

Our findings have important implications for organisations. First, based on our results, social image concerns will matter within organisations, where employees are usually

¹⁰The proportion of uninformed and informed agents who do not want to know their score is reported in Table [10](#).

well known to their employers. However, they may not matter as much in settings where the employee’s identity is less visible to the employer, e.g., when employees are recruited as freelancers on online platforms. Managers must also pay attention to the gender composition of the workforce as our results imply that social image concerns do not equally influence the behaviour of male and female employees. Second, to avoid suboptimal decisions driven by social image concerns, organisations could offer optional online training to all employees. Offering optional training to all means that employees who lack some skills will not have to ask for training and thus signal lower ability. Moreover, providing the training online allows employees to take the training without their superior and other employees necessarily knowing about it.

Some avenues for future research follow. First, future research could study whether people care about their social image because they feel pride when they are perceived as good or because they feel shame when they are perceived as bad.¹¹ Second, our paper opens the door to exploring the impact of social image concerns in strategic settings. We hope future work will investigate the role of social image concerns in other strategic environments where they may be relevant, such as employees working in teams.

¹¹In the context of our paper, one could vary the proportion of high-type agents to study whether the impact of social image concerns on project choice is driven by the pride an agent feels when perceived as a high type or by the shame they feel when perceived as a low type. See Appendix A (“Comparative statics”) for further details.

References

- Acemoglu, D., Autor, D., Hazell, J., and Restrepo, P. (2022). Artificial intelligence and jobs: Evidence from online vacancies. *Journal of Labor Economics*, 40(S1):S293–S340.
- Aghion, P. and Tirole, J. (1997). Formal and real authority in organizations. *Journal of political economy*, 105(1):1–29.
- Alonso, R. and Matouschek, N. (2008). Optimal delegation. *The Review of Economic Studies*, 75(1):259–293.
- Andreoni, J. and Bernheim, B. D. (2009). Social image and the 50–50 norm: A theoretical and experimental analysis of audience effects. *Econometrica*, 77(5):1607–1636.
- Ariely, D., Bracha, A., and Meier, S. (2009). Doing good or doing well? image motivation and monetary incentives in behaving prosocially. *The American Economic Review*, 99(1):544–555.
- Backus, P., Cubel, M., Guid, M., Sánchez-Pagés, S., and López Mañas, E. (2023). Gender, competition, and performance: Evidence from chess players. *Quantitative Economics*, 14(1):349–380.
- Bartling, B., Fehr, E., Maréchal, M. A., and Schunk, D. (2009). Egalitarianism and competitiveness. *American Economic Review*, 99(2):93–98.
- Bašić, Z. and Quercia, S. (2022). The influence of self and social image concerns on lying. *Games and Economic Behavior*, 133:162–169.
- Bénabou, R. and Tirole, J. (2006). Incentives and prosocial behavior. *American economic review*, 96(5):1652–1678.
- Bursztyn, L. and Jensen, R. (2015). How does peer pressure affect educational investments? *The quarterly journal of economics*, 130(3):1329–1367.
- Bursztyn, L. and Jensen, R. (2017). Social image and economic behavior in the field: Identifying, understanding and shaping social pressure. *Annual Review of Economics*, 9:131–153.
- Charness, G. and Gneezy, U. (2010). Portfolio choice and risk attitudes: An experiment. *Economic inquiry*, 48(1):133–146.
- Ellingsen, T. and Johannesson, M. (2008). Pride and prejudice: The human side of incentive theory. *American economic review*, 98(3):990–1008.

- Ewers, M. and Zimmermann, F. (2015). Image and misreporting. *Journal of the European Economic Association*, 13(2):363–380.
- Friedrichsen, J. and Engelmann, D. (2018). Who cares about social image? *European Economic Review*, 110:61–77.
- Friedrichsen, J., König, T., and Schmacker, R. (2018). Social image concerns and welfare take-up. *Journal of Public Economics*, 168:174–192.
- Gneezy, U., Niederle, M., and Rustichini, A. (2003). Performance in competitive environments: Gender differences. *The Quarterly Journal of Economics*, 118(3):1049–1074.
- Gneezy, U. and Rustichini, A. (2004). Gender and competition at a young age. *American Economic Review*, 94(2):377–381.
- Guo, J. and Recalde, M. P. (2023). Overriding in teams: The role of beliefs, social image, and gender. *Management Science*, 69(4):2239–2262.
- Katok, E. and Siems, E. (2011). Why genius leads to adversity: Experimental evidence on the reputational effects of task difficulty choices. *Management Science*, 57(6):1042–1054.
- Khalmetski, K. and Sliwka, D. (2019). Disguising lies-image concerns and partial lying in cheating games. *American Economic Journal: Microeconomics*, 11(4):79–110.
- Köszegi, B. (2006). Ego utility, overconfidence, and task choice. *Journal of the European Economic Association*, 4(4):673–707.
- Lacetera, N. and Macis, M. (2010). Do all material incentives for pro-social activities backfire? the response to cash and non-cash incentives for blood donations. *Journal of Economic Psychology*, 31(4):738–748.
- Murad, Z., Stavropoulou, C., and Cookson, G. (2019). Incentives and gender in a multi-task setting: An experimental study with real-effort tasks. *Plos one*, 14(3):e0213080.
- Najipoor-Schutte, K. and Patton, D. (2018). Survey: 68Harvard Business Review.
- Petrishcheva, V. (2023). Willful ignorance and reference dependence of self-image concerns.

Appendix

A Proofs

Proof of Remark 1

See that $a = \arg \max_x U_{H,0}(x)$ and $\emptyset = \arg \max_x U_{L,0}(x)$ because of Assumption (1). The action and its resulting outcome y fully separate types and the principal prefers to delegate because of Assumption (2) as $\gamma ps_P + (1 - \gamma)r_P > r_P$. $\square$

Proof of Lemma 1

Suppose by contradiction that there exists an H-type agent with $\eta' > 0$ who is indifferent between $x = a$ and $x = \emptyset$. For such agent, the following must hold:

$$r_A + \eta' \{ \tau(R)H + [1 - \tau(R)]L \} = ps_A + \eta' \{ [p\tau(S) + (1-p)\tau(F)]H + [p[1 - \tau(S)] + (1-p)[1 - \tau(F)]]L \}.$$

Since $ps_A > r_A$ by Assumption (1), it must be that

$$\eta' \{ [p\tau(S) + (1-p)\tau(F)]H + [p[1 - \tau(S)] + (1-p)[1 - \tau(F)]]L \} < \eta' \{ \tau(R)H + [1 - \tau(R)]L \},$$

which implies that all H-type agents with $\eta > \eta'$ will have a strict preference for $a = \emptyset$ and only $G(\eta')$ H-type agents will choose $x = a$. However, because $y = S$ (respectively, $y = F$) is more likely when $\theta = H$ (respectively, when $\theta = L$), the above equation also implies that

$$\tau(R)H + [1 - \tau(R)]L > [q\tau(S) + (1 - q)\tau(F)]H + [q[1 - \tau(S)] + (1 - q)[1 - \tau(F)]]L.$$

Since $r_A > qs_A$ because of Assumption (1), all L-type agents would choose $x = \emptyset$. Then, if an H-type agent chooses $x = a$, he would reveal that is a high-type, i.e., $\tau(S) = \tau(F) = 1$, and the above equation cannot hold. $\square$

Proof of Lemma 2

Suppose by contradiction that there exists such a fully-separating equilibrium, i.e., $Pr[\theta = H | y \in \{S, F\}] = 1$. Then low-type agents who are sufficiently concerned about their social image (i.e., low-type agents for whom $\eta > r_A - qs_A$) would prefer to deviate and choose $x = a$. It follows that $Pr[\theta = H | y \in \{S, F\}] < 1$. $\square$

Proof of Proposition 1

First notice that choosing $a = \emptyset$ reveals that $\theta = L$ because of Lemma 1. Then, the cutoff value of η for which an L-type agent is indifferent between $x = a$ and $x = \emptyset$, denoted by

$\hat{\eta}$, satisfies:

$$r_A + \hat{\eta}L = qs_A + \hat{\eta} \cdot \left\{ [q\tau(S) + (1-q)\tau(F)]H + [q[1-\tau(S)] + (1-q)[1-\tau(F)]]L \right\},$$

where

$$\begin{aligned} \tau(S) = Pr[\theta = H|y = S] &= \frac{Pr[y = S|\theta = H]Pr[\theta = H]}{Pr[y = S|\theta = H]Pr[\theta = H] + Pr[y = S|\theta = L]Pr[\theta = L]} \\ &= \frac{p\gamma}{p\gamma + q[1 - G(\hat{\eta})](1 - \gamma)}; \end{aligned}$$

and

$$\begin{aligned} \tau(F) = Pr[\theta = H|y = F] &= \frac{Pr[y = F|\theta = H]Pr[\theta = H]}{Pr[y = F|\theta = H]Pr[\theta = H] + Pr[y = F|\theta = L]Pr[\theta = L]} \\ &= \frac{(1-p)\gamma}{(1-p)\gamma + (1-q)[1 - G(\hat{\eta})](1 - \gamma)}. \end{aligned}$$

The indifference condition can be rewritten as:

$$r_A - qs_A = \hat{\eta} \left[q\tau(S) + (1-q)\tau(F) \right] H - \hat{\eta} \left[1 - q[1 - \tau(S)] - (1-q)[1 - \tau(F)] \right] L.$$

Note that the LHS is independent of η , whereas the RHS is strictly increasing in η , which means that there exists at most one admissible value of $\hat{\eta}$ that satisfies the above equality. That is, all L -type agents with $\eta \in [0, \hat{\eta}]$ strictly prefer $x = \emptyset$ to $x = a$ and all L -type agents with $\eta > \hat{\eta}$ strictly prefer $x = a$ to $x = \emptyset$. □

Proof of Proposition 2

Suppose that the principal believes that all H-type agents and a proportion $1 - G(\bar{\eta})$ of the L-type agents choose $x = a$. Then, she would be indifferent between choosing $d = 1$ and $d = 0$ if

$$\gamma ps_P + (1 - \gamma) \left[[1 - G(\bar{\eta})]qs_P + G(\bar{\eta})r_P \right] = r_P.$$

The above can be rewritten as:

$$\gamma(ps_P - r_P) = (1 - \gamma) \left[r_P - G(\bar{\eta})r_P - [1 - G(\bar{\eta})]qs_P \right],$$

and rearranging we obtain (4). As the principal rationally anticipates that $G(\hat{\eta})$ L-type agents choose $x = \emptyset$, the statement of the proposition follows. Lastly, see that $\bar{\eta}$ decreases when γ increases. □

Proof of Proposition 3

First, observe that the choice of H-type agents is unaffected by the parameter V . It is easy to see that $\frac{\partial \hat{\eta}}{\partial V} < 0$ and, therefore, it becomes more likely that $\hat{\eta} < \bar{\eta}$ when V increases. □

Comparative statics

In this subsection, we perform some comparative statics, studying the effect of different parameter values on $\hat{\eta}$.

It is easy to see that $\hat{\eta}$ is increasing with r_A and decreasing with s_A . The effect of q and p is ambiguous. The relationship with γ is more intricate. Let

$$Z := r_A - qs_A - \hat{\eta} \left[q\tau(S) + (1-q)\tau(F) \right] H + \hat{\eta} \left[1 - q[1 - \tau(S)] - (1-q)[1 - \tau(F)] \right] L = 0.$$

Applying the implicit function theorem, it holds that

$$\frac{\partial \hat{\eta}}{\partial \gamma} = -\frac{\frac{\partial H}{\partial \gamma}}{\frac{\partial H}{\partial \hat{\eta}}} = -\frac{-\hat{\eta} \left[q\frac{\partial \tau(S)}{\partial \gamma} + (1-q)\frac{\partial \tau(F)}{\partial \gamma} \right] (H-L) - \hat{\eta} \left[W_1 \frac{\delta H}{\delta \gamma} - W_2 \frac{\delta L}{\delta \gamma} \right]}{-W_1 H + W_2 L - \hat{\eta} \left[q\frac{\partial \tau(S)}{\partial \hat{\eta}} + (1-q)\frac{\partial \tau(F)}{\partial \hat{\eta}} \right] (H-L)},$$

where

$$W_1 := q\tau(S) + (1-q)\tau(F); \quad W_2 := q[1 - \tau(S)] + (1-q)[1 - \tau(F)].$$

Note that the denominator is negative because $\frac{\partial \tau(S)}{\partial \hat{\eta}} > 0$ and $\frac{\partial \tau(F)}{\partial \hat{\eta}} > 0$. Hence, the sign of $\frac{\partial \hat{\eta}}{\partial \gamma}$ coincides with the sign of the numerator. The first term is negative because $\frac{\partial \tau(S)}{\partial \gamma} > 0$ and $\frac{\partial \tau(F)}{\partial \gamma} > 0$, whereas the second term is ambiguous. The first term of the numerator captures the information value of a higher γ : the observation of $y = S$ or $y = F$ is more likely to come from an H-type agent as there is a higher fraction of such agents in the population. This unambiguously reduces the threshold $\hat{\eta}$. However, when γ increases the pride of being seen as one of the many high-types weakly diminishes, whereas the shame of being seen as one of the fewer low-types grows larger. Therefore, if people are primarily motivated by shame an increase in γ unambiguously reduces $\hat{\eta}$. Conversely, if individuals are primarily motivated by pride, the effect is ambiguous.

B Additional results

B.1 Derived loss across treatments

The payoff-maximising outcome is computed as follows:

$$\frac{1}{2}(0.75)(12 + 10) + \frac{1}{2}(4 + 6) = 13.25,$$

where informed agents choose the box that opens with a 75%, while uninformed agents select the outside option. Conversely, the joint payoff when principals refrain from delegating is simply 10. Accordingly, the potential gain from delegation is 3.25.

To compute the overall expected payoff under delegation in each treatment, we utilize the experimental results to determine the proportion of informed and uninformed agents

who select the box or the outside option. We consider the expected payoffs each agent-principal pair would have obtained if that agent's choice had been implemented. The corresponding percentages are then calculated as follows:

$$\frac{\text{Overall payoff in one treatment} - 10}{13.25 - 10}.$$

In the *Baseline* treatment, all informed agents and 48% of the uninformed agents choose the box. This implies that the expected joint payoff in this treatment is:

$$\frac{1}{2}(0.75)(12 + 10) + \frac{1}{2}[0.52(4 + 6) + 0.48(0.333)(12 + 10)] = 12.17.$$

The gain from delegation is

$$\frac{12.17 - 10}{13.25 - 10} = 0.668.$$

In the *Principal learns* treatment, 93% of the informed agents and 47% of the uninformed agents choose the box. This implies that the expected joint payoff in this treatment is:

$$\frac{1}{2}[(0.93)(0.75)(12 + 10) + 0.07(4 + 6)] + \frac{1}{2}[0.53(4 + 6) + 0.47(0.333)(12 + 10)] = 11.96$$

The gain from delegation is

$$\frac{11.96 - 10}{13.25 - 10} = 0.605.$$

In the *Photo* treatment, 90% of the informed agents and 64% of the uninformed agents choose the box. This implies that the expected joint payoff in this treatment is:

$$\frac{1}{2}[(0.9)(0.75)(12 + 10) + 0.1(4 + 6)] + \frac{1}{2}[0.36(4 + 6) + 0.64(0.333)(12 + 10)] = 11.48.$$

The gain from delegation is

$$\frac{11.48 - 10}{13.25 - 10} = 0.457.$$

B.2 Additional tables

Table 7: Impact of informing the principal on uninformed agents' choice

DV: Agent chooses box				
	(1)	(2)	(3)	(4)
Principal learns	0.357 (0.651)	0.323 (0.645)	0.345 (0.646)	0.173 (0.696)
Female	1.455** (0.683)	1.448** (0.684)	1.457** (0.683)	1.451** (0.690)
Principal learns $\times$ Female	-0.644 (1.049)	-0.567 (1.062)	-0.633 (1.044)	-0.670 (1.048)
Investment		0.111 (0.427)	0.171 (0.453)	0.215 (0.477)
Prosocial			-0.508 (0.650)	-0.562 (0.660)
See score				-0.992 (0.890)
People known				0.068 (0.131)
Constant	-0.762* (0.461)	-0.876 (0.693)	-0.830 (0.716)	-0.022 (1.103)
N_I	70	70	70	70

*** $p < 0.01$; ** $p < 0.05$; * $p < 0.1$

Logit regression. Standard errors in parentheses. N_I denotes the number of uninformed agents in the *Baseline* and the *Principal learns* treatment.

Table 8: Impact of lifting anonymity on uninformed agents' choice

DV: Agent chooses box				
	(1)	(2)	(3)	(4)
Not anonymous	1.475*** (0.555)	1.642*** (0.578)	1.691*** (0.615)	1.746*** (0.645)
Female	1.176** (0.511)	1.355** (0.548)	1.309** (0.549)	1.296** (0.550)
Not anonymous $\times$ Female	-1.753** (0.790)	-1.929** (0.806)	-2.009** (0.832)	-1.928** (0.851)
Investment		0.639** (0.317)	0.672** (0.328)	0.689** (0.335)
Prosocial			-0.872* (0.482)	-0.901* (0.480)
People known				0.033 (0.116)
See score				-1.241 (0.802)
Constant	-0.588* (0.323)	-1.362** (0.562)	-1.180** (0.562)	-0.125 (0.953)
N	120	120	120	120

*** $p < 0.01$; ** $p < 0.05$; * $p < 0.1$

Logit regression. Standard errors in parentheses. N_I denotes the number of uninformed agents in the three treatments.

Table 9: Marginal effect of lifting anonymity (aggregate)

DV: Agent chooses box				
	(1)	(2)	(3)	(4)
Not anonymous	0.169* (0.091)	0.178** (0.088)	0.173** (0.087)	0.175** (0.087)
N_I	120	120	120	120

*** $p < 0.01$; ** $p < 0.05$; * $p < 0.1$

Logit regression. Standard errors in parentheses. All coefficients are average marginal effects of showing each agent's photo to their principal on uninformed agents' probability of choosing a box. N_I denotes the number of uninformed agents in the three treatments.

Table 10: Frequency of agents choosing not to see their score

Do not see the score	Male	Female	Total
Uninformed agents	.08 (5)	.06 (3)	.07 (8)
Informed agents	.00 (0)	.3 (2)	.02 (2)

Number of observations in parentheses.

C Experimental Instructions

Screen 1

Welcome

Thank you for participating in this experiment.

Please do not communicate with other participants and do not use cell phones or other electronic devices at any time during the experiment.

If you have any questions, please **Raise your hand and an experimenter will come to you** to answer your question privately.

For your participation in this experiment, you will receive a **€5 participation fee**. You can also win **additional money depending on your decisions** in the experiment and the decisions of other participants.

You will receive your payment by bank transfer. No other participant will be informed of your payment, and you will not be informed of any other participant's payment.

Next

Screen 2A

[only shown to participants A]

Your Role: A

Each participant has been randomly assigned to one of two possible roles: A or B. You have been assigned to role A. Like all other participants in role A, you will now complete a 30-question IQ test, which is frequently used to measure intelligence.

You should try to answer as many questions as possible. You will have up to 10 minutes to complete the test. Your final payment may depend on your performance in the test compared to the performance of other participants. More information will be provided after the test.

You are paired with a participant who has been assigned to Role B. No participant in Role B can complete the test; however, they will be able to see it.

[Next](#)

Screen 2B

[only shown to participants B]

Your Role: B

Each participant has been randomly assigned to one of two possible roles: A or B. You have been assigned to role B. All participants in role A will now complete a 30-question IQ test, which is frequently used to measure intelligence.

Like all participants in role B, you will not complete the test, but you can see the test as participants in role A complete it. Your final payment may depend on how some of the A participants linked to you perform in the test relative to other participants. More information will be provided after the test.

You are paired with four participants in Role A:

- Participant A2
- Participant A3
- Participant A4
- Participant A5

[Next](#)

Screen 3

[RAVEN'S TEST]

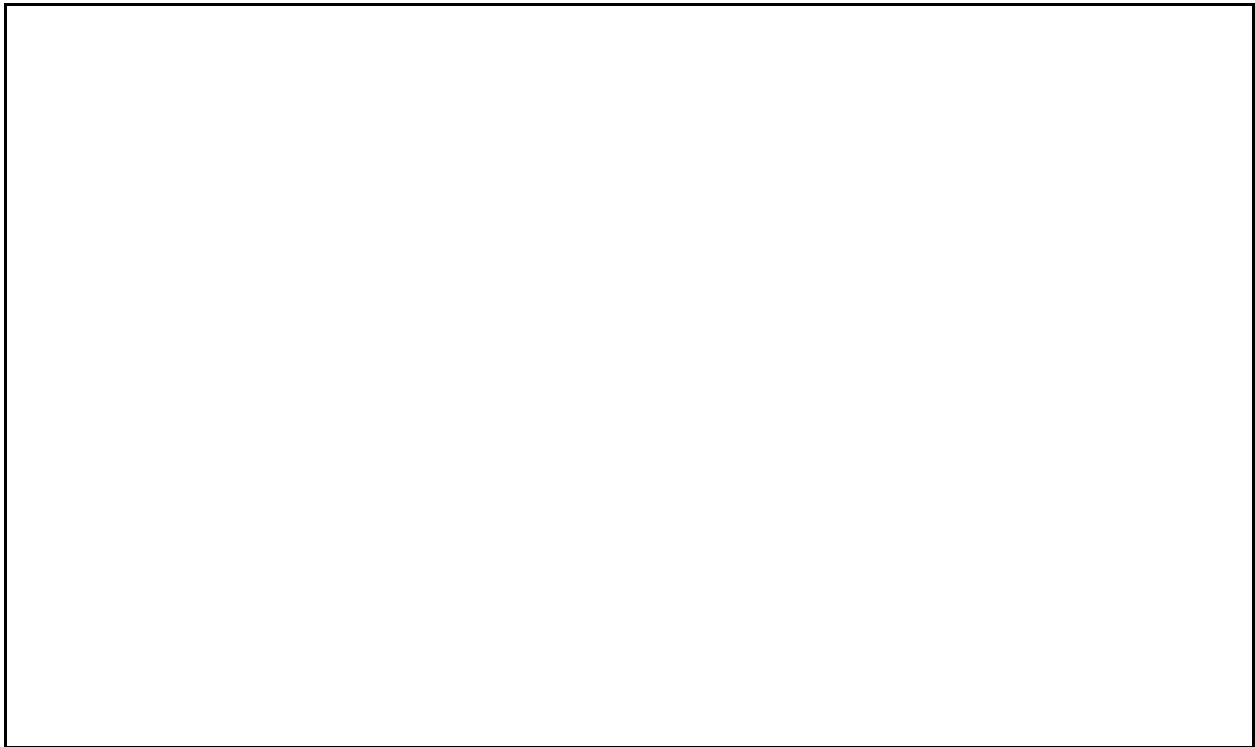
Screen 4

[ONLY PHOTO TREATMENT]

Photograph

As part of the experiment, all participants must take a photo. Depending on their role in the experiment, a participant's photo may be shown to members of their group (more information will be provided later). Please note that the experimenters will not see your photo and it will be deleted after the experiment.

Please take a clear photograph, facing the camera.



Capture Image

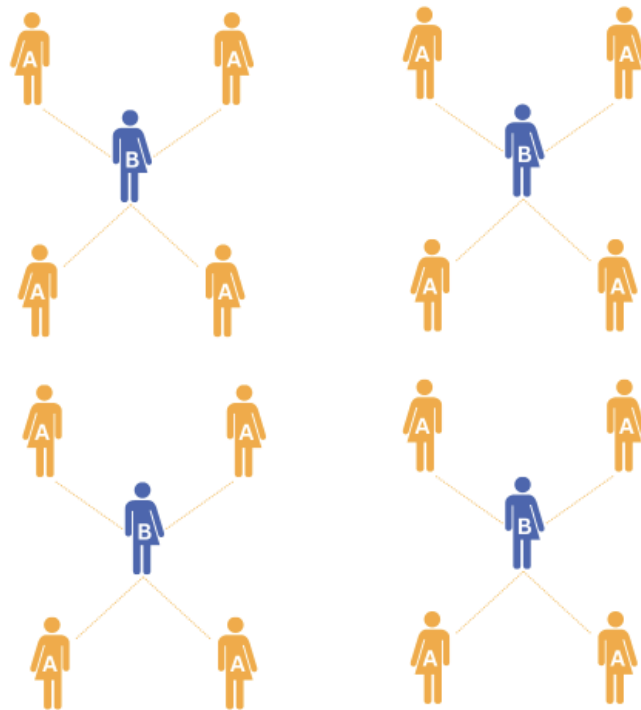
Upload Image

Screen 5

Please pay full attention to these instructions. At the end of the instructions, you will be asked some comprehension questions.

Groups:

Each participant in role B (hereinafter referred to as “participant B”) is linked to 4 participants in role A (hereinafter referred to as “participants A”), as shown below:



Please note that, although groups always contain five participants, the total number of groups in each session may vary depending on how many people participate. Each session includes between 1 and 4 groups.

Initial situation:

There are three boxes: the **blue** box, the **green** box, and the **red** box. One of these three boxes contains money: €**12** for participant A and €**10** for participant B. The other two are empty. If one of these two boxes is selected, both participants get 0 euros.

With a probability of **25%** The box containing the money is locked and cannot be opened. This means that with a probability of **25%** the box containing the money will leave both participants with €**0**.

If no box is chosen (a participant “abstains”), participant B wins €**6** and participant A wins €**4**.

At first, no participant knows which of the three boxes contains the money. They know that each box has the same probability of containing the money.

Note: Which box contains the money may differ between participants A.



Only the participants A who scored in the top 50% in the test among the participants A in the classroom will find out which box contains the money, while the others will not find out. In case of ties, these will be resolved based on the time of completion of the test.

Participants B never find out which box contains the money

[BASELINE AND PHOTO TREATMENTS]: , and they never know which participants A scored above 50% in the test score among the participants A in the classroom.

[PRINCIPAL LEARNS TREATMENT], but they know which of their Participant A scored above 50% in the test among the Participants A in the classroom.

There will be three stages:

Stage 1: Each B participant must decide whether to abstain or allow the participants A to whom he is linked (hereinafter “her/his participants A”) to make a decision.

If a participant B decides to allow her/his participants A to make a decision, one of them will be selected at random and her/his decision will determine the payout for participant B.

Stage 2: Each participant A learns if s/he scored in the top 50% in the test. If s/he is in the top 50% in the test, s/he immediately finds out which box contains the money.

If s/he is not in the top 50% in the test, s/he does not find out.

Stage 3 [BASELINE TREATMENT]: Without knowing the decision made by their linked participant B in Stage 1, each participant A must decide whether they want to abstain or choose a box.

Her/his Participant B will observe this decision and the resulting payment.

Specifically, for each of her/his participants A, a B participant will observe one of the following possible consequences:

- A chooses a box and B receives **€10**
- A chooses a box and B receives **€0**
- A abstains and B receives **€6**

Note: The decision made by each participant A will only be implemented if her/his participant B has decided to allow them to make a choice in Stage 1.

If her/his participant B has decided to abstain, the decision made by participants A will have no consequences.

Note: All interactions between a Participant B and her/his four Participants A will be anonymous.

Stage 3 [PRINCIPAL LEARNS TREATMENT]: Without knowing the decision made by their linked participant B in Stage 1, each participant A must decide whether they want to abstain or choose a box.

Not only will her/his participant B observe this decision and the resulting payment, but also whether her/his participant A knows which box contains the money.

Specifically, for each of his/her participants A, a B participant will observe one of the following possible consequences:

In the case in which A discovers which box contains the money:

- A scored above 50% in the test, chooses a box and B receives € 10
- A scored above 50% in the test, chooses a box and B receives € 0
- A scored above 50% in the test, abstains and B receives € 6

In the case in which A does not discover which box contains the money:

- A did not score above 50% in the test, chooses a box and B receives € 10
- A did not score above 50% in the test, chooses a box and B receives € 0
- A did not score above 50% in the test, abstains and B receives € 6

Note: The decision made by each participant A will only be implemented if her/his participant B has decided to allow them to make a choice in Stage 1.

If her/his participant B has decided to abstain, the decision made by participants A will have no consequences.

Note: All interactions between a Participant B and her/his four Participants A will be anonymous.

Stage 3 [PHOTO TREATMENT]: Without knowing the decision made by their linked participant B in Stage 1, each participant A must decide whether they want to abstain or choose a box.

His/her Participant B will observe this decision, the resulting payment, and her/his photograph.

Specifically, for each of his/her participants A, a B participant will observe one of the following possible consequences:



A[number] chose a box and B receives €10



A[number] chose a box and B receives €0



A[number] abstained and B receives €6

Note: The decision made by each participant A will only be implemented if her/his participant B has decided to allow them to make a choice in Stage 1.

If her/his participant B has decided to abstain, the decision made by participants A will have no consequences.

Next

Screen 6

Comprehension questions

Do participants in role B find out which participants in role A scored in the top 50% in the test?

No

Yes

[BASELINE TREATMENT]: What do the participants in role B know about each of their participants in role A?

If they opened a box and the resulting payment.

Their score in the test.

Their true identities

[PRINCIPAL LEARNS TREATMENT]: What do the participants in role B know about each of their participants in role A?

If they scored in the top 50% in the test and therefore find out which box contains the money, whether they opened a box, and the resulting payment.

Their score in the test.

Their true identities

[PHOTO TREATMENT] What do the participants in role B know about each of their participants in role A?

If they opened a box and the resulting payment.

Their score in the test.

If they opened a box, the resulting payment and they also see their photograph.

Do participants in Role B discover which box contains the money?

No

If they scored in the top 50% in the test.

Yes, always

Can participants in Role A choose a box even if they don't know which box contains the money?

No

Yes

How much does a participant in role B earn when s/he allows her/his participants in role A to make a decision and choose a box?

€12

€10

It depends on whether the box contains the money

It depends on whether the box contains the money, whether the box is locked, and which participant in role A is selected for payment.

How much does a participant in role A earn when her/his participant in role B allows her/him to make a decision and the participant A chooses a box?

€4

€10

It depends on whether the box contains the money

It depends on whether the box contains the money and whether the box is locked.

How much does a participant in role A earn when s/he abstains?

€4

€6

€0

How much does a participant in role B earn when s/he abstains?

€4

€6

€0

[Submit answers](#)

Screen 7

Error checking

Screen 8 [only PHOTO TREATMENT]

Estos son sus participantes A:



Participante A1



Participante A2



Participante A3



Participante A4

Next

Screen 9A1

You have not scored in the top 50%, so you won't find out which box contains the money.

[Next](#)

Screen 9A2

Congratulations, you have scored in the top 50%!

The box that contains the money is the **red** box (*example*).

[Next](#)

Screen 10A

Decision for Participants A

Please make a decision.

- I choose the red box.
- I choose the blue box.
- I choose the green box.
- I abstain from choosing a box.

[See Instructions](#)

Screen 10B

Decision for Participants in Role B.

- Allow participants A to make a decision.
- Abstain

[See Instructions](#)

Screen 11

[BASELINE TREATMENT]

Here is the information about your participants' decisions in Role A and your own payment (*example*)

- Participant A1 decided to choose a box and you got a payment of 0 Euros.
- Participant A2 decided to choose a box and you got a payment of 0 Euros.
- Participant A4 decided to choose a box and you got a payment of 10 Euros.
- Participant A5 decided to choose a box and you got a payment of 10 Euros. **This participant was chosen at random to determine your payment.**

Your payment for this part is: €10.00

Screen 11

[PHOTO TREATMENT]

[Photo treatment is the same as the baseline treatment, with the exception that the Participant A's picture is displayed next to that Participant's choice and the resulting payment - as shown below:]

Información para el Participante en el Rol B

Aquí está la información sobre las decisiones de tus participantes en el Rol A y tus propios pagos:

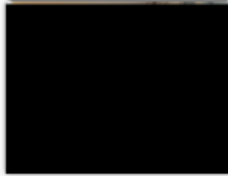


Participante A2 decidió elegir una caja y tu obtuviste un pago de 0 Euros.

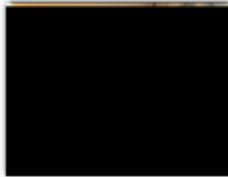


Participante A3 decidió elegir una caja y tu obtuviste un pago de 10 Euros.

Este participante fue elegido al azar para determinar tu ganancia.



Participante A4 decidió elegir una caja y tu obtuviste un pago de 0 Euros.



Participante A5 decidió elegir una caja y tu obtuviste un pago de 0 Euros.

Tu pago de esta parte es: €10.00 Euros

Next

Screen 11

[PRINCIPAL LEARNS TREATMENT]

Here is the information about your participants' decisions in Role A and your own payment (*example*)

- Participant A1 did not score above 50% in the test, s/he decided to choose a box and you got a payment of 10 Euros. **This participant was chosen at random to determine your payment.**
- Participant A3 scored above 50% in the test, s/he decided to choose a box and you got a payment of 10 Euros.
- Participant A4 scored above 50% in the test, s/he decided to choose a box and you got a payment of 0 Euros.
- Participant A5 did not score above 50% in the test, s/he decided to choose a box and you got a payment of 10 Euros.

Your payment for this part is: €10.00

Screen 12

(For Participants A ONLY)

You can decide whether you want to see your percentage of correct answers in the test and your ranking in relation to other participants, or continue without seeing this information.

Show my score and ranking

Proceed without showing my results

Screen 13

Test Results

(Example) Your test score: 0%

(Example) Your ranking: #2 out of 16 participants.

[Next](#)

Screen 14

You have been given €2 and must choose the portion of this amount (between €0 and €2, included) that you wish to invest.

The amount that is not invested is yours. Any amount invested will be multiplied by 2.5 with 50% probability and will become 0 with 50% probability.

Only one participant in this room will be chosen randomly and paid for this decision.

How much would you like to invest?

€

Amounts are rounded to two decimal places.

[Next](#)

Screen 15

In the following table, there are 4 rows corresponding to Decision 1, Decision 2, Decision 3 and Decision 4. In each decision, you must choose between two options: Left Option (L) and Right Option (R). In each decision, the option chosen will specify a win for you and a win for another participant in the room. Everyone in the room will make the same four decisions.

Once all participants in the room have made their decisions, one of the participants will be randomly chosen. For the participant who is chosen, one of their four decisions will be randomly selected and payments will be made based on the option they have chosen in that decision.

In addition, another participant in the room will be randomly selected to receive the amount of € that corresponds to “the other participant.”

Please select an option (Option L or Option R) for each decision:

Decisión	Opción L	Opción R
Decisión 1	☐ 2€ para ti 2€ para el otro participante	☐ 1€ para ti 1€ para el otro participante
Decisión 2	☐ 2€ para ti 2€ para el otro participante	☐ 3€ para ti 1€ para el otro participante
Decisión 3	☐ 2€ para ti 2€ para el otro participante	☐ 2€ para ti 4€ para el otro participante
Decisión 4	☐ 2€ para ti 2€ para el otro participante	☐ 3€ para ti 5€ para el otro participante

Siguiente

Screen 16

Are you a student?

Yes

No

How many participants do you know in this classroom?

[ONLY FOR Players B in PHOTO TREATMENT] How many participants do you know in your group?

Are you an undergraduate or graduate student?

Undergraduate

Postgraduate

Which year of studies are you in?

What is your field of study?

Business Administration

Economy

Other

Are you employed?

Yes

No

What country are you from?

How old are you?

What is your gender?

Man

Women

Other

[Submit answers](#)

Screen 17

Final Payment Information

Your final payment will be: €19.00 (*Example*)

Please enter your name for payment processing. This information will not be associated with your decisions in this experiment and will only be handled by administrative staff.

[Submit and End Experiment](#)