

Relational Influence

Heikki Rantakari*

University of Rochester

Simon Business School

April 3, 2017

Abstract

An uninformed principal elicits non-contractible recommendations from a privately informed agent regarding the quality of projects. The agent is biased in favor of implementation and no credible communication is possible in a one-shot setting. In a repeated setting, the fear of losing future influence can sustain informative communication, but the agent's willingness to remain truthful depends on the extent to which he expects the principal to listen to him. In a stationary equilibrium, the principal always implements mediocre projects at a sub-optimally high frequency to reward honesty, while she may either favor or discriminate against high-quality projects. In a non-stationary equilibrium, the principal will further condition the agent's future influence on today's proposals, with the admission of mediocre alternatives rewarded with increased future influence while rejections of high-quality projects are further punished by lowering the agent's future influence. The acceptance of high-quality projects builds up influence when the agent's current influence is not too high, but erodes the influence when the agent is already highly influential.

*contact: heikki.rantakari@simon.rochester.edu. I would like to thank Michael Raith, Luis Rayo, Alex Wolinsky and the participants of the fall 2016 NBER Organizational Economics Meeting for helpful comments. The usual disclaimer applies.

1 Introduction

Issues of decision-making under strategic information transmission have been increasingly recognized to be of crucial importance for organizational performance. As noted by Cyert and March (1963 [1992]), "[w]here different parts of the organization have responsibility for different pieces of information..., [we would expect] attempts to manipulate information as a device for manipulating the decision." (p.79). Following the theoretical frameworks introduced by Crawford and Sobel (1982) and Milgrom and Roberts (1988), a large and growing literature has examined how incentive conflicts lead to attempts at manipulating information, the resulting loss of information and how to manage such losses, leading to the analysis of issues of delegation, mediation, information management and beyond.

A particular aspect of such relationships, especially in organizational settings, is that they are typically ongoing. For example, a company does not just choose a single R&D project, build a single factory or choose where to locate a new logistics center. Instead, the organization faces an ongoing sequence of such decisions, and it is the same group of organizational members that are involved in the decision-making process. This ongoing nature of the decision-making process then allows the development of relationships, so that even if the quality of current recommendations cannot be verified on the spot, the parties can learn about the quality of past recommendations from the outcomes that have resulted.

This paper constructs a simple model of such relationships and considers how the parties can use the history of the relationship as a basis for current behavior and to sustain a relationship that is better for both parties than a one-shot interaction. In the setting, a principal needs to decide whether to implement a project. The project can be either mediocre or good, information which is learned only the agent. The agent makes a recommendation to the principal regarding the quality of the project, after which the cost of implementation is publicly observed. The cost of implementation may be a literal cost or reflect the random evolution of other actions the principal might take and is informed about. The principal wants to implement the project only when its value exceeds the cost of implementation, while the agent cares only about the value of the project. If the principal chooses to implement the project, the value of the project is learned before the next choice is made, while if the principal chooses against implementation, the value is not learned. This asymmetry captures the idea that we learn more about recommendations that are followed relative to recommendations that are not followed.

To focus on how the principal can manage the relationship without money, I assume that monetary transfers are not possible. Instead, the object of interest is the principal's decision rule, which determines when the project is implemented, conditional on whether the agent submits either a weak or strong report regarding the quality of the project and the realized cost of implementation. Since the decision rule effectively determines how likely it is that the agent's proposal is implemented, it determines the agent's influence in the relationship. The goal of the analysis is to examine how the principal can use this relational influence of the agent to manage the relationship.

The first basic observation is that trust creates value because information is valuable to both parties. Thus, when the agent is sufficiently patient, he is willing to truthfully reveal the quality of

the project even when the principal follows her preferred decision rule (henceforth first-best). If the agent misleads the principal, the principal may learn that after the fact and will stop trusting the agent's recommendations going forward, leading to less informed decision-making and worse payoffs to both parties.

When the agent is not patient enough, then the temptation to push for the acceptance of mediocre projects becomes too high under the first-best decision rule. The question is then what the principal can do to maintain the relationship. I begin by considering a stationary equilibrium (the current decision rule is not a function of past history), where the problem consists of choosing the optimal distortions in the implementation rule. Noting that the main constraint that we need to satisfy is keeping the agent honest about mediocre projects, the basic distortions are two-fold. First, the principal will always bias the acceptance rule in favor of mediocre alternatives and thus implements some mediocre projects at a loss. Such favoritism is beneficial because it simultaneously increases the value of the relationship to the agent and lowers the agent's temptation to exaggerate. Second, the principal will either favor or discriminate against high-quality projects. On one hand, favoring high-quality projects increases the agent's ongoing influence and thus makes him less willing to sacrifice that influence for an immediate gain. On the other hand, favoring high-quality projects also increases the gains from exaggeration by making the strong proposal more influential relative to a weak proposal. When the agent is sufficiently patient, the first effect (future value) dominates, while when the agent is less patient, the second effect (higher immediate gain) dominates. Thus, an agent of intermediate patience is rewarded with higher than the first-best level of influence (with the principal over-implementing both high- and low-quality projects), while an agent with low patience has lower than the first-best level of influence (the discrimination against high-quality proposals more than outweighs the remaining over-implementation of mediocre projects).

Finally, while both types of projects may be implemented with excessive frequency, the relative bias in favor of mediocre projects is growing in the level of impatience. In practical terms, the equilibrium thus provides a simple explanation for corporate socialism, whereby the internal allocation of resources is less responsive to differences in profitability than suggested by a simple NPV criterion. Relatedly, variation in the patience of the agent leads to either over- or under-investment relative to the first-best use of resources.

Allowing the current decision rule to depend on past outcomes considerably enriches the play of the game. While an explicit solution for the optimal strategies is currently beyond reach, we can establish three salient features that such an equilibrium will contain. First, the principal will reward the agent for the admission of mediocre projects by increasing his future influence. The benefits of this delayed reward are two-fold. First, because the cost of distorting decisions is convex in the distortion, spreading the reward for honesty over several periods is more cost-effective than settling up immediately through an excessive likelihood of implementing a mediocre project. Second, because the agent also values the implementation of high-quality projects more, a promise of increased future influence is attractive to the agent – instead of settling up now by implementing a mediocre project, the principal promises to give the agent more favorable treatment in the future, where the agent may have a high-quality project available. This result provides a simple rationale for basic quid pro quo arrangements, where honesty today is rewarded by favorable treatment in the future. For

example, a department admitting that their favorite job candidate is mediocre is promised priority in the hiring process next year, or a division admitting for limited investment opportunities today is promised easier access to funding for any new projects the following year.

Second, in addition to rewarding the admission of mediocre alternatives, the principal will lower the agent's future influence if she chooses to reject a project receiving strong support. The reason for this result is that when a proposal is not accepted, its quality is not learned. This lack of information limits the punishment that the principal can impose on the agent when he attempts to mislead her. By lowering the agent's payoff when the proposal is rejected, the incentives to exaggerate are decreased. However, because even honest recommendations of high-quality projects are sometimes rejected (and, indeed, in equilibrium, only honest recommendations arise), terminating the relationship is too harsh of a punishment. Instead, the principal responds by lowering the future influence of the agent without fully stopping trusting him. In practice, this feature resembles a situation where an agent falls out of favor with the principal – once a strong proposal is rejected, the agent's odds of getting future projects through are lowered.

Third, while rejections of high-quality projects are followed by reductions in future influence, the acceptance of high-quality projects is followed by either increases or decreases in the future influence. When the agent's current level of influence is not too high (including the principal's preferred level of influence), then the acceptance of a high-quality project is followed by an increase in future influence. The reason is that such a delayed reward is simply another means deferred compensation, limiting the need for immediate settling up through the current implementation decisions. When the current level of influence is high, however, then the acceptance of high-quality projects can be followed by reductions in future influence because additional promises become increasingly expensive to the principal. In other words, the agent can first build up influence through the acceptance of high-quality projects, but for higher levels, the agent starts to cash in on that influence when his proposals are implemented.

Finally, the challenge created by the distortions in the stage game, whether in the stationary or non-stationary equilibrium, is that as the agent becomes increasingly impatient, the distortions needed to keep the agent truthful grow. Then, the need for the principal to honor the promises made in equilibrium can put a cap on both the distortions that can be sustained in the stage game and any promises of additional future influence. In such cases, the relationship may become unsustainable, with the distortions needed to keep the agent truthful being too large to be credible for the principal to follow. Then, no informative equilibrium will exist.

The rest of the manuscript is organized as follows. Section 2 discusses the related literature and section 3 outlines the model. Section 4 provides a preliminary examination of the framework, illustrating the full set of payoffs attainable in the stage game and the conditions under which the first-best equilibrium can be obtained. Section 5 derives the optimal stationary equilibrium, and Section 6 considers dynamics. Section 7 concludes and discusses some potential extensions. Appendix B illustrates a game with a continuous project space for the agent to illustrate how the stationary decision rule is optimally determined in such settings.

2 Related Literature

This paper lies in the intersection of the literatures on repeated games and strategic communication. The five main papers existing in this intersection are Alonso and Matouschek (2008), Kolotilin and Li (2015), Campbell (2015), Li et al. (2015) and Lipnowski and Ramos (2015). The first two papers consider the classic repeated-game setting of full ex post observability of outcomes. Alonso and Matouschek (2008) consider the Crawford and Sobel (1982) setting with a long-lived principal interacting with a sequence of myopic agents. The value of the ongoing relationship helps the principal to choose a decision closer to the agent's preferences, facilitating communication. If the principal is sufficiently patient, she achieves the optimal delegation set (and thus the maximum payoff to the principal in the absence of transfers). Kolotilin and Li (2015) extend the CS setting to a long-lived agent and the ability to make transfers. Transfers make the communication problem trivial by having the agent signal his information with an associated transfer, and the issue is how to manage decision-making by the principal who underweighs the agent's payoff in her favored decision. I consider a qualitatively different decision problem, which introduces the asymmetric learning regarding the quality of recommendations based on whether they are followed or not.

The remaining three papers consider settings that are qualitatively closer to the present model, with the agent making recommendations regarding which (if any) projects to implement, but consider the opposite extreme of no learning of the outcomes. The strategies can thus be based only on the observed history of recommendations. In Campbell (2015), the agent uses his relational capital to recommend a project as long as it is good enough, which uses his relational capital until it is exhausted. The capital is never replenished. Lipnowski and Ramos (2015) is closest to the present paper, where the relational capital is both replenished and used over time, but where the replenishment occurs when the agent recommends rejection while the capital is used whenever the agent recommends acceptance.¹ Finally, in Li et al. (2015) the agent can recommend either his ideal project or a project that is better for the principal, but that project may not be available. When the agent recommends his own project, the continuation value must punish the agent to keep him honest, drifting the equilibrium towards not listening to him, while recommending the project that is better for the principal increases the continuation value, with an increased likelihood that the agent gets to choose his preferred project whenever he wants to. The key differences of the present paper in relation to these three contributions are as follows. First, in all three papers the question is simply whether the principal follows the agent's recommendation, thus not allowing for the richer manipulation of the acceptance rule to manage the relationship, which is the focus here. Second, the assumptions regarding observability are different, with my setting allowing for partial observability of the outcomes. The key resulting difference in dynamics is that "acceptance" always depletes capital while "rejection" always builds up capital in the above papers, while the present setting allows the agent to build up capital both through the admission of mediocre projects and, initially, the acceptance of high-quality projects, while capital is consumed by the rejection of high-quality projects and, when the current stock of influence is high enough, the acceptance of high-quality

¹For a related paper, see Guo and Horner (2015).

projects.

More broadly, the present paper relates to the large literature on repeated games with private information that has followed Abreu et al. (1990), where the focus on the use of continuation values and distortions in the behavior (instead of monetary transfers) to sustain the equilibrium is present in, e.g. Athey and Bagwell (2001) and Athey et al. (2004) on colluding with private information, Mobius (2001) and Hauser and Hopenhayn (2008) on favor trading, Li and Matouschek (2008) on enforcing the payment of bonuses by the principal, Padro i Miquel and Yared (2012) on managing the moral hazard problem of an intermediary in maintaining the rule of law, and Andrews and Barron (2014) on managing multiple supply relationships, just to mention a few. The building up and using of influence is analogous to the favor trading literature following Mobius (2001), albeit in a different context, while the use of rejection of high-quality alternatives to punish the agent is analogous to the punishment mechanism in Li and Matouschek (2008) and Padro i Miquel and Yared (2012).

In terms of strategic communication, the paper considers a variant of the framework analyzed in Li et al. (2016), Rantakari (2016), Garfagnini et al. (2014) and Chakraborty and Yilmaz (2013), among others, where a decision-maker needs to choose among discrete alternatives, based on the recommendation(s) of an agent or multiple agents. The setting retains the discrete nature of the final choice, but introduces the continuity of private information for the principal to smooth out the decision problem.

3 Model

I consider a repeated advisory relationship between an agent and a principal. In the stage game, the agent has access to a "project," the value of which is given by $\theta_i \in \{\theta_L, \theta_H\}$, with $0 \leq \theta_L < \theta_H \leq 1$. Let the probability of the high-quality alternative be given by p . The agent observes privately the value of the project, and makes a recommendation $m_i \in \{m_L, m_H\}$ to the principal as to the quality of the project (since the project quality is binary, we can restrict our attention to binary messages). The recommendations are soft information (cheap talk), and the principal interprets the message according to equilibrium play to form beliefs regarding the quality of the project, $E(\theta|m_i)$.

Following the recommendation, the principal's outside option, c , is drawn from a known distribution F , which, for tractability, I assume to be $U[0, 1]$. Once the outside option is realized and publicly observed, the principal chooses whether to adopt the project of the agent or choose the outside option. For concreteness, we can take the value to be the expected revenue generated by a given project, while the outside option is the cost of investment. Then, if the principal accepts the agent's project, the payoff is given by $\theta_i - c$ while the outside option is no investment, with normalized payoff of 0. The agent, on the other hand, does not bear any of the cost of investment, and the agent's payoff is given by θ_i if the principal invests and 0 if the principal doesn't invest. The principal observes her payoff at the end of the period, so she will learn whether the agent told the truth if she follows the recommendation, but does not if she chooses the outside option. The discount rates are $\delta_A, \delta_P < 1$ for the agent and the principal, respectively. The projects and costs are distributed iid over time, with each period involving new draws for both θ_i and c .

No transfers: If the parties had access to (unbounded) transfers, the solution would be trivial even in the static setting, as the agent could signal his private information through voluntary transfers. In many settings, however, transfers are either not available or are limited for various reasons, including risks of collusion and rent-seeking activities. Thus, I make the opposite assumption, where no transfers are available between the agent and the principal. Instead, the relationship will be sustained by the principal's decision rule, $\Pr(A|m_i, c)$, which specifies the (A)ccceptance probability following the agent's message and the commonly observed principal's state. As noted by Cyert and March (1963), "Side payments, far from being the incidental distribution of a fixed, transferable booty, represent the central process of goal specification. That is, a significant number of these payments are in the form of policy commitments." (p.35)

Other assumptions: To maintain the tractability of the analysis and to be able to explore the dynamics of the relationship, I make a number of further simplifying assumptions. To mention a few, I assume a binary state for the private signal, publicly observable cost of investment, single agent, and perfect observability of the outcome when a project is implemented. The qualitative logic of the analysis remains if the agent's state is continuous, but the decision rule becomes naturally richer, highlighting differences among low- and medium-quality projects. This setting is discussed in Appendix B. The publicly observable principal's state allows us to focus on managing the relationship with only that one agent. An interesting avenue for future work is the examination of how to manage the relationship when multiple agents hold relevant information to the decision. Finally, a valuable extension would be to consider the imperfect observability of the outcomes, either by introducing noise into the principal's payoff or by allowing for imperfect information for the agent. However, the set of assumptions allows us to focus on how the principal can both instantaneously and dynamically manage the relationship with the agent when the only tool available is the decision rule, $\Pr(A|m_i, c)$. Finally, I assume that the parties are engaged in the relationship whether the principal trusts the agent or not, so that the common threat point is playing the uninformative equilibrium.

4 Preliminaries - Feasible payoffs and First-Best

Before considering the equilibrium of the model, I will first consider the stage-game payoffs and the basic tradeoffs involved. This will help to summarize the structure of the model and thus provide insight into the results that follow later.

4.1 Payoff structure and the payoff possibilities frontier

Consider first the expected payoff of the principal and the agent. Assuming truth-telling by the agent, we can write their payoffs as

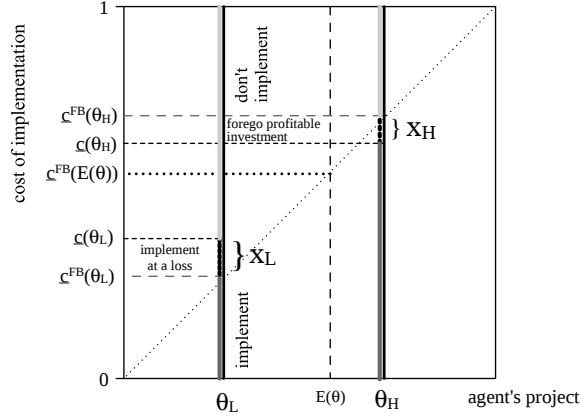


Figure 1: Illustrating the principal's decision rule

$$v_P = \sum_{i \in \{L, H\}} \Pr(\theta_i) \int_c \Pr(A|c, \theta_i) (\theta_i - c) dF(c) \quad \text{and} \quad v_A = \sum_{i \in \{L, H\}} \Pr(\theta_i) \int_c \Pr(A|c, \theta_i) \theta_i dF(c),$$

where $\Pr(A|c, \theta_i)$ indicates the probability of acceptance. As the first preliminary observation, note that since the agent cares only about implementation, any reasonable acceptance rule will take a threshold struture, where the principal accepts the project as long as her cost of implementation is below some threshold, $\underline{c}(\theta_i)$. Given the assumption that the costs are uniformly distributed, we can then write the expected payoffs as

$$v_P = \sum_{i \in \{L, H\}} \Pr(\theta_i) \underline{c}(\theta_i) \left(\theta_i - \frac{\underline{c}(\theta_i)}{2} \right) \quad \text{and} \quad v_A = \sum_{i \in \{L, H\}} \Pr(\theta_i) \underline{c}(\theta_i) \theta_i.$$

Second, note that the first-best threshold for implementation is $\underline{c}^{FB}(\theta_i) = \theta_i$. Thus, we can write any implementation rule simply as $\underline{c}(\theta_i) = \theta_i + x_i$, where x_i is the distortion away from the first-best decision rule. The principal's strategy can thus be summarized by the pair $\{x_L, x_H\}$, the examination of which will be the focus of the analysis. If, on the other hand, the principal makes no use of the available information, then the expected quality of the project is $E(\theta) = p\theta_H + (1-p)\theta_L$ and the optimal (common) threshold for implementation is then $\underline{c}^{FB}(E(\theta)) = E(\theta)$. In this case, the expected payoffs become simply

$$\underline{v}_P = \frac{E(\theta)^2}{2} \quad \text{and} \quad \underline{v}_A = E(\theta)^2.$$

The basic structure is illustrated in Figure 1. In this illustration, the principal is biasing her decision rule in favor of the agent when the agent is making the weaker recommendation, while discriminating against the agent when he makes the strong recommendation ($x_L > 0, x_H < 0$).

To measure the value of the agent's information, we can solve for $(v_P - \underline{v}_P)$ and $(v_A - \underline{v}_A)$, which are given by

$$u_P = (v_P - \underline{v}_P) = \frac{1}{2} \left[(1-p)p(\theta_H - \theta_L)^2 - px_H^2 - (1-p)x_L^2 \right] \quad (1)$$

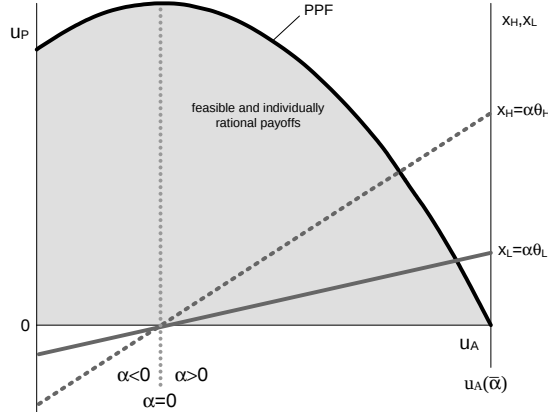


Figure 2: Feasible payoffs

$$u_A = (v_A - \underline{v}_A) = p(1-p)(\theta_H - \theta_L)^2 + p\theta_H x_H + (1-p)\theta_L x_L \quad (2)$$

The value of the agent's information to the principal (and to himself) is thus proportional to $\phi(\theta) = (1-p)p(\theta_H - \theta_L)^2$. The value is lowered for the principal, however, whenever the decision rule is distorted away from the first-best, so that $x_i \neq 0$. The cost of distortions is naturally convex in the size of the distortions, with the loss given by $px_H^2 + (1-p)x_L^2$. The agent, on the other hand, benefits from a more favorable decision rule, where the value to the agent is given by $p\theta_H x_H + (1-p)\theta_L x_L$. I will call this component the relational influence of the agent, as any increase in x_i makes the principal more likely to follow the suggestion of the agent and thus increase his payoff. As a note on notation going forward, v_i, V_i will be used to denote player i 's gross payoff in the stage-game and the resulting net present value, while $u_i = v_i - \underline{v}_i$ and $U_i = V_i - \frac{1}{1-\delta_i}\underline{v}_i$ denote the net value generated by the information revealed.

The second question is what is the overall set of feasible payoffs. To this end, we must solve for the payoff possibilities frontier, consisting of maximizing u_P conditional on delivering a given value u_A to the agent. This involves solving for the least-cost deviation needed to deliver a given value to the agent. From equations 1 and 2 we obtain the following Lemma:

Lemma 1 Optimal distortions: *distortions (x_L, x_H) are on the payoff frontier if and only if $\frac{\theta_L}{\theta_H} = \frac{x_L}{x_H}$. We can thus characterize the frontier with a single coefficient α , where $x_i = \alpha\theta_i$.*

Proof. Holding the agent's expected payoff constant gives the tradeoff between the distortions as $\frac{dx_H}{dx_L} = -\frac{(1-p)\theta_L}{p\theta_H}$ while maximizing the principal's payoff requires $px_H \frac{dx_H}{dx_L} + (1-p)x_L = 0$. Together, these give $\frac{\theta_L}{\theta_H} = \frac{x_L}{x_H}$, which means any efficient distortion must satisfy $x_i = \alpha\theta_i$. ■

An important converse of this lemma is that if the decision rule is not proportional, then we are bounded away from the frontier. To finish characterizing the frontier, define $\bar{\alpha} = \sqrt{\frac{\phi(\theta)}{p\theta_H^2 + (1-p)\theta_L^2}} < 1$ as the maximum level of influence that can be given to the agent in any game, defined by $u_P(\bar{\alpha}) = 0$. Finally, note that (i) $u_P(\alpha = 0)$ gives the principal's preferred equilibrium, (ii) $\frac{du_P(\alpha)}{du_A(\alpha)} = |\alpha|$ and $u_A(\alpha)$ is proportional to α , so that the frontier is concave, and that (iii) each player is guaranteed a net payoff at least as high as zero. The resulting frontier together with all feasible payoffs is illustrated in Figure 2, together with an illustration for the decision rule distortions that attain the boundary.²

4.2 Obtaining the first-best (and other parts of the frontier)

Consider now whether obtaining the principal's preferred equilibrium (first-best) is feasible. Since the maximizer for the principal is unique (at $\alpha = 0$), it can be obtained only if it is self-generating (as is any other point on the frontier due to its strict concavity). Given that this is the principal's preferred equilibrium, we do not need to worry about her incentive-compatibility constraint. The only relevant constraint is for the agent to make the correct recommendation, which in turn can be binding only when faced with a mediocre alternative.

By telling the truth, the agent guarantees himself an expected payoff of

$$\Pr(A|m_L)\theta_L + \delta_A V_A. \quad (3)$$

In other words, his proposal is accepted with probability $\Pr(A|m_L) = \theta_L$, and whether or not the proposal is accepted, the principal continues to trust the agent, keeping the game on the same continuation path, with value V_A . In contrast, if he chooses to deviate, his expected payoff is given by

$$\Pr(A|m_H)\theta_L + \delta_A [\Pr(A|m_H)V_A^{dev} + \Pr(R|m_H)V_A]. \quad (4)$$

He thus increases the immediate acceptance probability to $\Pr(A|m_H) = \theta_H$ by exaggerating the quality of the proposal, but now the lie is detected with probability θ_H , in which case the principal stops trusting the agent. However, if the proposal is still rejected, the principal learns nothing and thus the game remains on the equilibrium path, with value V_A . The strategies thus constitute an equilibrium if and only if

$$(\Pr(A|m_H) - \Pr(A|m_L))\theta_L \leq \delta_A \Pr(A|m_H) [V_A - V_A^{dev}], \quad (5)$$

which, in the case of $\alpha = 0$, simplifies to

² $\theta_H = 3\theta_L$

$$\frac{\theta_L (\Pr(A|m_H) - \Pr(A|m_L))}{\theta_L (\Pr(A|m_H) - \Pr(A|m_L)) + \Pr(A|m_H)\phi(\theta)} \leq \delta_A. \quad (6)$$

Thus, as long as the agent is sufficiently patient, the principal is able to obtain her preferred outcome, simply by utilizing her preferred decision rule and stopping to trust the agent whenever he is caught lying.

The more interesting case arises when the condition is not satisfied, and so the first-best outcome is no longer obtainable by the principal. The broader analysis is undertaken below, and I will only make two preliminary observations regarding the payoff frontier. First, if the first-best is not obtainable at a given (δ_A, δ_P) , then no point with $\alpha < 0$ is self-generating since it provides a lower payoff to both the agent and the principal. Second, while the first-best may not be attainable at a given (δ_A, δ_P) , points with $\alpha > 0$ may be self-generating. In particular, letting $\eta(\theta) = p\theta_H^2 + (1-p)\theta_L^2$ denote the rate of value transfer as we alter α , then a point on the frontier will be attainable as long as

$$\delta_P \geq \frac{2|\alpha|\theta_H}{2|\alpha|\theta_H + [\phi(\theta) - \alpha^2\eta(\theta)]} \quad \text{and} \quad \delta_A \geq \frac{(\theta_H - \theta_L)\theta_L}{(\theta_H - \theta_L)\theta_L + \theta_H[\phi(\theta) + \alpha\eta(\theta)]}. \quad (7)$$

In other words, as we increase α , while the equilibrium becomes more sustainable for the agent because of the higher value of the relationship (note that a proportional increase in x_H and x_L leaves the reneging temptation unchanged), it becomes less likely to be incentive-compatible for the principal because sustaining the equilibrium requires increasingly large distortions in the decision rule. Once both parties are sufficiently impatient, no point on the frontier can be attained.

5 Stationary equilibrium

Having considered the basic structure of the problem, we can now consider the repeated game itself. In this section, I will consider the optimal stationary equilibrium for the principal, where the distortions (x_L, x_H) are independent of the history of the play. In the next section, I will consider how the principal can do better by considering history-dependent strategies and what are the basic tradeoffs involved.

The distortions need to satisfy two incentive-compatibility constraints. First, as above, truth-telling needs to be in the agent's self-interest. We can write this constraint as

$$\begin{aligned} \Pr(A|m_L)\theta_L + \frac{\delta_A}{1-\delta_A}v_A(x_L, x_H) &\geq \Pr(A|m_H)\theta_L + \frac{\delta_A}{1-\delta_A}[\Pr(A|m_H)v_A + (1 - \Pr(A|m_H))v_A(x_L, x_H)] \\ &\Leftrightarrow \\ \frac{\delta_A}{1-\delta_A}u_A(x_L, x_H) &\geq \frac{(\Pr(A|m_H) - \Pr(A|m_L))}{\Pr(A|m_H)}\theta_L, \end{aligned} \quad (8)$$

where $\Pr(A|m_i) = \theta_i + x_i$, with x_i as the distortion in the decision rule. This expression contains the main insights regarding the agent's truth-telling constraint. First, the constraint can be binding only for the lower-quality alternative, as the gain comes from increasing the probability of acceptance. Second, even if the agent has a temptation to misrepresent only the low-quality alternative, the principal may optimally alter her decision rule for both alternatives.

Because the temptation to lie arises from the incremental increase in the acceptance probability, $\frac{\Pr(A|m_H) - \Pr(A|m_L)}{\Pr(A|m_H)}$, the first means through which the principal will manage the constraint is to increase the acceptance probability when the agent sends a weak recommendation. This increase in $\Pr(A|m_L)$ will both increase the agent's continuation value $u_A(x_L, x_H)$ and relax the reneging temptation. The second means is through altering the acceptance probability following a strong recommendation, $\Pr(A|m_H)$. Here, however, the effects go in opposite directions: increasing the acceptance probability following a strong recommendation both increases the continuation value (relaxing the constraint) and increases the immediate gain to deviation (tightening the constraint). As a result, as we will see below, the equilibrium distortion may be in either direction.

For the principal, the incentive-compatibility constraint arises from the fact that by deviating from the decision rule, she is able to save the distortion x_i . Thus, for her, the reneging temptation arises for both recommendations as long as $x_i \neq 0$. This constraint can be written as

$$\max |x_i| \leq \frac{\delta_P}{1 - \delta_P} u_P(x_H, x_L), \quad (9)$$

We can thus write the principal's maximization problem as

$$\begin{aligned} \min_{x_H, x_L} & \quad (px_H^2 + (1-p)x_L^2) \\ \text{s.t.} & \quad \frac{\delta_A}{1 - \delta_A} u_A(x_L, x_H) \geq \frac{(\theta_H + x_H - \theta_L - x_L)}{(\theta_H + x_H)} \theta_L \\ & \quad \max |x_i| \leq \frac{\delta_P}{1 - \delta_P} u_P(x_H, x_L). \end{aligned}$$

I will consider the solution in two steps. First, I will ignore the principal's IC constraint and consider the solution when we only need to satisfy the agent's truth-telling constraint and later I will re-introduce the principal's IC constraint. The logic behind the solution is easiest to illustrate graphically, as done in Figure 3. The principal's indifference curves are ellipses, as denoted by the dotted lines, with utility improving towards the origin. The agent's truth-telling constraints are denoted by the rotating lines with the dashed arrows denoting the rotation as the agent becomes increasingly impatient, where incentive-compatibility requires that the distortions lie to the north-east (or south-east) of the relevant truth-telling constraint.

As we saw above, when the agent is sufficiently patient, the principal is able to achieve her desired outcome at $x_H^{FB} = x_L^{FB} = 0$. As the agent becomes sufficiently impatient, however, this solution is no longer feasible. Instead, the principal chooses the distortion that obtains the lowest indifference curve, subject to satisfying the agent's IC constraint, and tracing this tangency point gives us the equilibrium (x_H^*, x_L^*) for any patience level by the agent. To understand the logic behind the path,

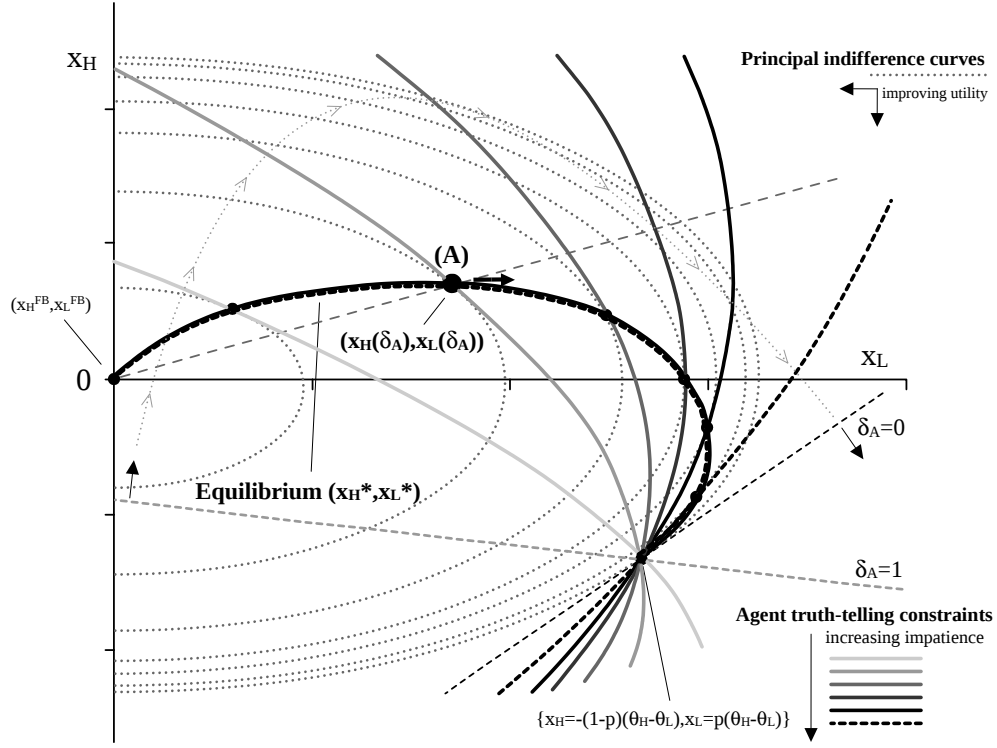


Figure 3: Deriving the optimal distortions

recall from above that increasing x_L decreases the reneging temptation while increasing continuation value, while increasing x_H increases both. When the agent is patient enough, the continuation value effect dominates the use of x_H , and thus the principal initially increases both x_H and x_L to sustain truthful communication. As δ_A increases, however, the relative usefulness of x_H decreases because of its growing relative impact on the reneging temptation. This effect is reflected in the "fanning out" of the truth-telling constraints (to maintain indifference, any reduction in x_L must be matched with an increasingly large increase in x_H), and eventually it becomes optimal for the principal to start decreasing x_H while still increasing x_L . When the agent becomes sufficiently impatient, it becomes optimal to start to discriminate against high-quality projects and x_H becomes negative. Finally, once the agent is very impatient, instead of further increasing x_L it becomes optimal to start shrinking both x_H and x_L , until the solution converges to $x_L = p(\theta_H - \theta_L)$, $x_H = -(1-p)(\theta_H - \theta_L)$ when the agent becomes fully myopic and no use of information is possible.³ This logic is formalized in the following proposition:

Proposition 2 *The optimal stationary decision rule:*

- (i) *The relative bias for low-quality projects is increasing in the impatience of the agent, with the ratio always bounded away from the efficient distortion: $\frac{x_H^*}{x_L^*} < \frac{\theta_H}{\theta_L}$ and $-\frac{d(\frac{x_H}{x_L})}{d\delta_A} < 0$.*

³Note that at this point, $\theta_H + x_H = \theta_L + x_L = E(\theta)$. Also, the truth-telling constraints rotate around this point since it is always incentive-compatible for the agent to tell the truth, no matter what the patience level, when the decision rule makes no use of that information.

(ii) There exists a range of (low but positive) patience levels for which the probability of implementing low-quality projects exceeds the implementation probability under no information: $\Pr(A|m_L, \delta_A) > E(\theta)$

Proof. See Appendix A.1 ■

The first part of the proposition formalizes the graphical representation of the solution, with the additional detail that because the ratio of distortions is bounded away from the efficient distortion and is everywhere decreasing, the solution is always represented by a clockwise rotation of the rays from the origin tracing through the tangency points, with the original slope strictly less than $\frac{\theta_H}{\theta_L}$. The second part highlights that the bias in favor of mediocre projects will for some patience levels be so high that the implementation probability exceeds the uninformed implementation threshold. The corollaries that follow from this result are as follows:

Corollary 3 Implications of the optimal stationary decision rule:

- (i) The relative probability of acceptance for low-quality projects is monotone increasing in the impatience of the agent: $-\frac{d(\frac{\Pr(A|L)}{\Pr(A|H)})}{d\delta_A} > 0$.
- (ii) The agent's payoff is initially increasing but later decreasing in his impatience, with a sufficiently impatient agent having a payoff below the first-best level: $u_A(x_H^*(\delta_A), x_L^*(\delta_A)) \geq u_A(0, 0)$ for $\delta_A \geq \underline{\delta}_A$ and $u_A(x_H^*(\delta_A), x_L^*(\delta_A)) < u_A(0, 0)$ for $\delta_A < \underline{\delta}_A$.
- (iii) The principal initially over-invests but later under-invests in expectation relative to the first-best: $E(c|x_H^*(\delta_A), x_L^*(\delta_A)) \geq E(c|0, 0)$ for $\delta_A \geq \tilde{\delta}_A$ and $E(c|x_H^*(\delta_A), x_L^*(\delta_A)) < E(c|0, 0)$ for $\delta_A < \tilde{\delta}_A$.

Part (i) of the corollary thus formalizes the result of corporate socialism, where the worse-quality projects are implemented disproportionately too often relative to the high-quality projects. The second part formalizes the behavior of relational influence, where the agent is incentivized with above-FB levels of influence for intermediate patience levels to sustain truth-telling but penalized with below-FB levels of influence when he is very impatient to avoid the abuse of influence. The third part makes the related observation in terms of expected investment costs, whereby the above-FB influence of the agent leads to excessive investment in the agent's projects, whereas for low patience levels, the discrimination against high-quality projects leads to below-FB levels of investment due to the poor ability to condition those investments on the actual quality of the projects.

Adding the principal's IC constraint: So far, the analysis only considered the impact of the agent's truth-telling constraint. As the agent becomes increasingly impatient, the distortions needed to maintain truth-telling grow in size and the principal may become tempted to deviate herself, putting a cap on the distortions that can be introduced to the decision rule. A more complete discussion of the constraint is provided in Appendix A.1, but the logic of the basic impact is illustrated in Figure 4. The key feature of the IC constraints is that each constraint can be bounded

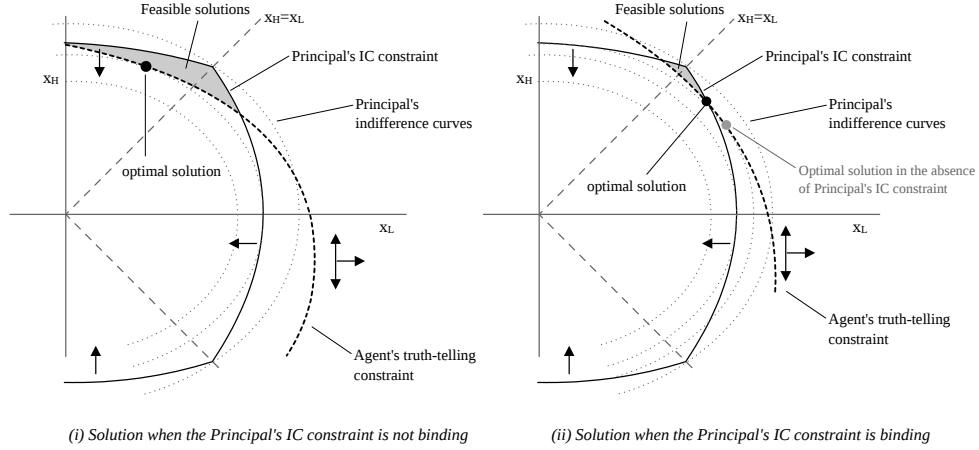


Figure 4: Illustrating the impact of the principal's IC constraint.

between the indifference curve that it touches on the diagonal and a rectangle that is drawn from that point. The basic logic is that, starting from the diagonal, decreasing x_L or x_H increases the continuation value to the principal and thus increases the maximal distortion that can be sustained for the other distortion. This creates the "bowing-out" effect. But because the principal discounts the future, the extent to which the constraint is relaxed depends on the discount rate. As $\delta_P \rightarrow 0$, the benefit disappears and the IC constraint converges to a rectangle (at the origin), while as $\delta_P \rightarrow 1$, the discounting disappears and the IC constraint converges to the matching (outermost) indifference curve.

Bringing the two together is then straightforward. Satisfying the agent's truth-telling constraint requires the distortions to be sufficiently large, while the principal's IC constraint requires them to be sufficiently small. Together they define the feasible set of solutions, within which the principal chooses the one that attains the lowest loss.

When the agent is sufficiently patient, then the feasible set contains the tangency point between the principal's indifference curve and the agent's truth-telling constraint. In this case, the principal attains her (constrained) preferred outcome and her IC constraint is irrelevant. Such a solution is illustrated in panel (i). But when the agent becomes less patient, the feasible set shrinks and the tangency point will lie outside the feasible set. Then, the principal will choose the best possible solution within the feasible set, which involves both constraints being just satisfied. As drawn in panel (ii), this involves the principal reducing x_L to make the decision rule credible for her to promise, while increasing x_H to restore truth-telling for the agent. As we increase the impatience of the agent further, the feasible set disappears and informative communication becomes impossible.

While panel (ii) illustrates the distortion only for one configuration of the parameters, the general logic extends to the other configurations, and is as follows. When the agent is patient enough, the principal's IC constraint is irrelevant. Once the principal's IC constraint becomes binding, it becomes generically binding only for one of the distortions x_i . Then, the principal can use the slack in the other IC constraint to continue to satisfy the agent's truth-telling constraint. As a result, once the principal's IC constraint becomes binding, the solution travels along the constraint towards

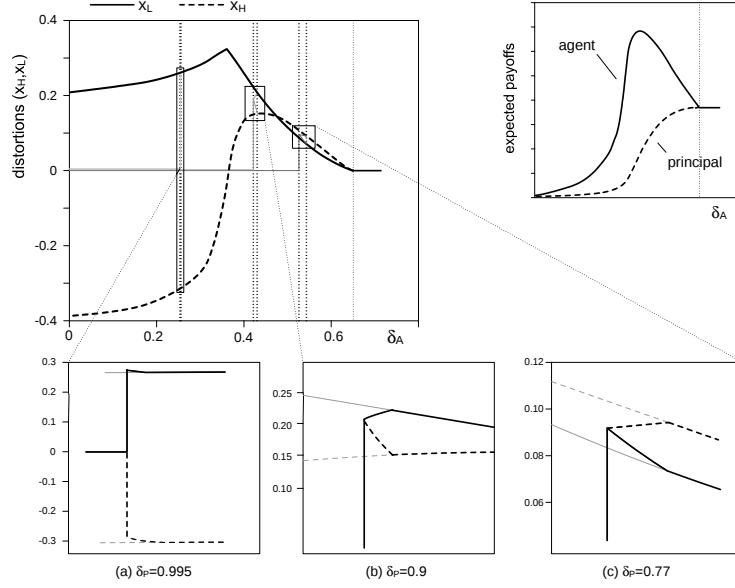


Figure 5: Optimal distortions and expected payoffs under stationary equilibrium - $\theta_L = 0.2, \theta_H = 0.8$ and $p = 0.35$.

the diagonal. Once on the diagonal, $|x_H| = |x_L|$, the principal's IC constraint is binding for both distortions and no further adjustment is possible. Then, any additional impatience by the agent collapses the relationship.

An example: A numeric illustration of the above discussion is illustrated in Figure 5. The top panel plots the distortions and the expected payoffs to the agent and the principal, illustrating the non-monotone behavior of the distortions and the resulting non-monotone payoff of the agent, under a perfectly patient principal.

The bottom of the figure (panels (a)-(c)) illustrate the effects of the principal's IC constraint at various patience levels. Naturally, the more impatient the principal, the earlier her IC constraint binds and the configuration of the levels of distortions at which the constraint becomes binding changes. But in each case, the same adjustment towards the diagonal occurs, until the equilibrium no longer exists. As a final observation, note that the agent may be locally better off due to these additional distortions before being hurt by the unsustainability of the relationship.

To summarize, the basic insights from the analysis of the stationary equilibrium are as follows. First, to limit the incentives to exaggerate, the optimal decision rule always exhibits corporate socialism, in the sense that mediocre projects have a relatively higher likelihood of acceptance than high-quality projects than the first-best outcome. Indeed, the acceptance probability of the mediocre projects may exceed the uninformed threshold. Second, the overall relative influence of the agent is non-monotone in his patience. This non-monotonicity followed from the result that increasing the likelihood of accepting high-quality projects both increased the continuation value to the agent (relaxing the truth-telling constraint) and increased the immediate gain from misleading

the principal (tightening the truth-telling constraint). When the agent is sufficiently patient, the first effect dominates and it is optimal to reward the agent with above-first-best relational influence to maintain truth-telling, whereas when the agent becomes more impatient, the temptation for the agent to abuse that high level of influence becomes too high and it becomes optimal to limit the agent's influence below the first-best level. In terms of project financing, this implies that the overall level of investment may be both above and below the first-best level.

6 Non-stationary strategies

While the distortions in the current influence of the agent can be used to sustain the relationship, a stationary policy fails to take into account the possibility of using changes in the future influence of the agent to achieve the same. The basic logic is relatively simple. Recall that the main constraint that the principal needs to worry about is to induce the agent to admit that his project is mediocre. We can increase the attractiveness of this admission by rewarding the agent with additional future influence following the admission of mediocre projects, whether they are accepted or not. For example, the dean of faculty may not hire a given job candidate or the CEO may not choose to finance a given project if given only lukewarm support by the department or division, but promises priority treatment in the future. This use of future influence brings two main benefits. First, it allows the principal to smooth out the reward for truth-telling over multiple periods: instead of settling up immediately through a higher acceptance probability of a mediocre project, the reward is spread over time by giving the agent higher expected level of influence in the future. Second, it brings about an important efficiency gain: since the agent values the implementation of high-quality projects more than low-quality projects, partially delaying the reward is valuable. Instead of having a low-quality alternative implemented today, he receives a promise of favorable treatment tomorrow, when he may have a high-quality project available. Finally, the reward of higher influence increases the agent's continuation value going forward, helping to sustain informative communication in the future as well.

In addition to rewarding the agent for the admission of mediocre projects, the principal can also lower the agent's influence if she chooses to reject a strong proposal by the agent. The reason is that when the agent makes a strong recommendation but the proposal is rejected, the true quality of that project is never learned and this limits the cost of deviation to the agent. To counter this, it will be optimal to follow such a rejected recommendation with a decrease in future influence. This result is akin to the punishment phase in games of imperfectly observed actions, such as the triggering of price wars in Green and Porter (1984) and lower effort levels by the agent in Li and Matouschek (2013).

To see these features in more detail, index the current influence of the agent by ω , with influence increasing in ω . Then, following the equilibrium strategy (truth-telling) gives the agent an expected payoff of

$$\Pr(A|L, \omega) \theta_L + \delta_A [\Pr(A|L, \omega) V_A(A|L, \omega) + (1 - \Pr(A|L, \omega)) V_A(R|L, \omega)], \quad (10)$$

where L indicates the lower recommendation and A indicates an acceptance (and R rejection). Similarly, if the agent lies, then his continuation payoff is given by

$$\Pr(A|H, \omega) \theta_L + \delta_A [\Pr(A|H, \omega) V_A^{dev} + (1 - \Pr(A|H, \omega)) V_A(R|H, \omega)] . \quad (11)$$

We can then combine these two equations to yield the new truth-telling constraint as

$$\Delta U_A(L, \omega) - \Pr(R|H, \omega) \Delta U_A(R|H, \omega) + \Pr(A|H, \omega) U_A(\omega) \geq \frac{1}{\delta_A} (\Pr(A|H, \omega) - \Pr(A|L, \omega)) \theta_L, \quad (12)$$

where $\Delta U_A(L, \omega)$ is the expected change in the continuation value following the admission of a mediocre alternative, $\Delta U_A(R|H, \omega)$ is the change in the continuation value following the rejection of a high-quality proposal, and $U_A(\omega)$ is the net present value in the current state before the project quality is realized.⁴ Note that if $\Delta U_A(L, \omega) = \Delta U_A(R|H, \omega) = 0$, we are back to the stationary equilibrium, while both $\Delta U_A(L, \omega) > 0$ and $\Delta U_A(R|H, \omega) < 0$ can be used to relax the constraint.

The second constraint that needs to be satisfied is the principal's promise-keeping constraint, which states that the expected payoff promised to the agent as a result of the outcome of the previous period, $U_A(\omega)$, must equal the expected payoff the agent will receive from that period stage game and the resulting continuation payoffs. In other words, we have

$$U_A(\omega) = \frac{u_A(\omega) + \delta_A (1 - p) \Delta U_A(L, \omega) + \delta_A p [\Pr(R|H, \omega) \Delta U_A(R|H, \omega) + \Pr(A|H, \omega) \Delta U_A(A|H, \omega)]}{(1 - \delta_A)}. \quad (13)$$

For the principal, the constraints are similar to before, with the additional effect of changes in the continuation value that can be used to incentivize the principal as well. In short, the constraints are now (noting that $x_i > 0$ implies that it is the acceptance decision that the principal would like to deviate from and vice versa):

$$\begin{aligned} x_i &\leq \delta_P (U_P(\omega) + \Delta U_P(A|i, \omega)) & \text{if } x_i > 0 \\ |x_i| &\leq \delta_P (U_P(\omega) + \Delta U_P(R|i, \omega)) & \text{if } x_i < 0 \end{aligned} . \quad (14)$$

Then, the principle of optimality implies that the principal wants to maximize, in each period, his current net value

$$U_P(\omega) = \frac{u_P(\omega) + \delta_P (1 - p) \Delta U_P(L, \omega) + \delta_P p [\Pr(R|H, \omega) \Delta U_P(R|H, \omega) + \Pr(A|H, \omega) \Delta U_P(A|H, \omega)]}{(1 - \delta_P)}, \quad (15)$$

⁴That is, $\Delta U_A(L, \omega) = \Pr(A|L, \omega) U_A(A|L, \omega) + (1 - \Pr(A|L, \omega)) U_A(R|L, \omega) - U_A(\omega)$ and $\Delta U_A(R|H, \omega) = U_A(R|H, \omega) - U_A(\omega)$.

subject to the agent's truth-telling and promise-keeping constraints (12 and 13) and her own incentive-compatibility constraints (14) over the current-period distortions ($\{x_{i,\omega}\}$) and promised changes in continuation values ($\{\Delta U_A(j|i, \omega)\}$), together with the initial promise of influence $U_A(\omega_0)$ in the beginning of the first period.

Now, to fully characterize the optimal solution, we would need to solve for the self-generating payoff set. Unfortunately, the richness of the principal's action space makes characterizing the set challenging.⁵ We can, however, use the principle of optimality to characterize the basic tradeoffs involved to obtain the economic logic behind the dynamics, which is performed in the next subsection. Having considered the basic tradeoffs and some features of the solution, I then provide a numeric solution to a simplified three-state variant of the model. For what follows, I assume that $\delta_A = \delta_P$ so that the strategic timing of payments to utilize time preferences will not play a role.

6.1 Characterizing the dynamics

To consider the dynamics of the optimal relationship, we can make two preliminary observations. First, because there is only one-sided asymmetric information, the optimal contract can utilize payoffs on the payoff frontier. No joint punishment inside the frontier is needed. Second, because the stage-game losses are convex in the magnitude of the distortions, the payoff frontier should be strictly concave even when considering the optimal dynamic strategies.⁶

Given these two observations, we can then consider the potential frontiers and the movements along them. But before considering the dynamics, it is instructive to consider the set of attainable payoffs under stationary strategies and the distortions associated with them. These are illustrated in Figure 6. Panel (a) illustrates the case of moderate patience. Recall that the frontier is obtained by the proportional distortion, $x_i = \alpha\theta_i$. For moderate patience, intermediate α can be sustained as a part of an equilibrium, and the frontier contains a segment of the feasible frontier. For lower levels of agent utility, the relationship is not valuable enough to the agent and the truth-telling constraint becomes binding. Then, to limit the gains to misleading the principal, we need to increase x_L and, to satisfy the target payoff for the agent, lower x_H . Eventually, however, these distortions become too big to be incentive-compatible to the principal and delivering a lower payoff to the agent is not possible in a stationary equilibrium.⁷ For higher levels of agent utility, it is the principal's constraint that becomes binding, and to satisfy that constraint we need to lower the favoritism towards high-quality projects and, to sustain the agent's payoff, increase the favoritism towards low-quality projects. Once $x_L \rightarrow x_H$, no further promises are possible and no higher payoff to the agent can be delivered. Panel (b) illustrates the same for a lower level of patience. Now, the parties are impatient enough that the frontier can no longer be attained. For lower agent utilities, the agent's truth-telling constraint is binding, forcing us to distort x_L above and x_H below the

⁵The concavity of the frontier implies that the optimal solution does not have the bang-bang property that would allow us to characterize the Pareto frontier simply through its extreme points.

⁶The convexity of the losses implies that any randomization between two stage game outcomes is strictly dominated by a deterministic move to an intermediate outcome.

⁷Note that random termination does not help since it would also punish the principal, making even the current distortions unsustainable.

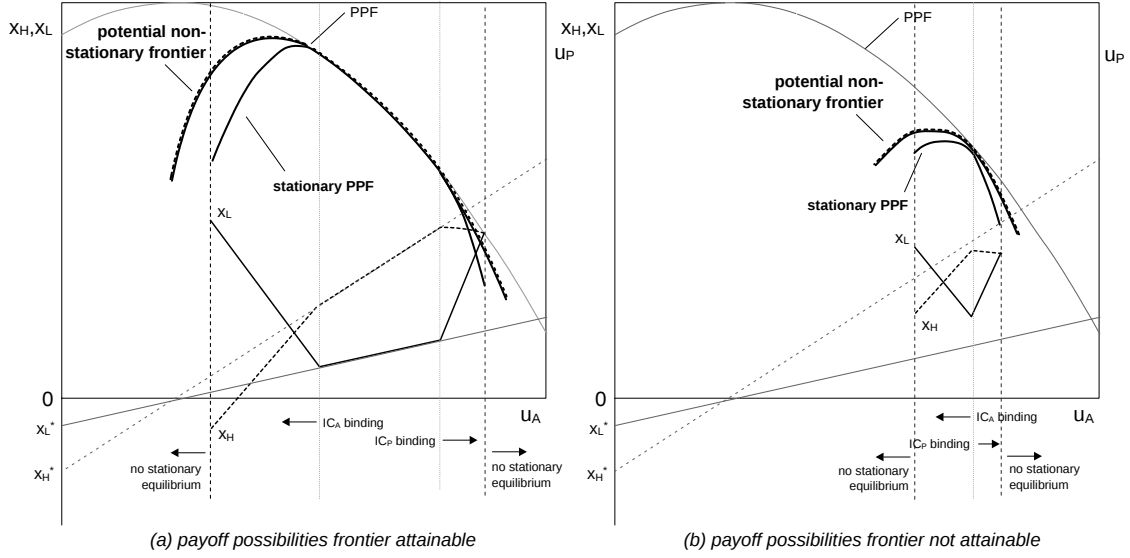


Figure 6: Illustrating the payoff frontiers

efficient distortion, while for higher agent utilities, the principal's IC binds and limits the difference in distortions.

For the frontier of the non-stationary payoffs, we know that they must be bounded between the feasible frontier and the stationary frontier. In addition, we can show that, with the exception of when the stationary equilibrium attains the feasible frontier, a non-stationary strategy can always do better for the principal. In particular, there is no absorbing state at either end of the non-stationary frontier, as given by the following lemma:

Lemma 4 *Consider the end-points of the non-stationary frontier where either the agent's truth-telling constraint and/or the principal's IC constraint binds. Then, there always exists a non-stationary strategy that dominates the static repetition of that state.*

Proof. See Appendix A.2 ■

The simple intuition behind this result follows from the constraints (12 and 14). When the agent's truth-telling constraint is binding in a low state, increasing the reward through the increase in future influence allows us to lower the current distortion x_L to deliver the desired continuation value to the agent in a more efficient manner. As an additional benefit, it relaxes the principal's IC constraint. Similarly, when the principal's IC constraint is binding in a high state, we can relax that constraint by rewarding the principal following the acceptance of a high-quality project by lowering the agent's future influence. While this reward lowers the agent's continuation value, it relaxes the principal's IC constraint enough that she can deliver the same continuation value through more efficient distortions in the high state. As a result, stationary repetition of the end-points will not be optimal, and the same logic applies to all points where either one of the constraints is binding. As

a result, the non-stationary frontier will always lie strictly above the stationary frontier whenever the stationary frontier does not reach the feasible frontier, and, importantly, extends the range of attainable payoffs for the agent outside the bounds of the stationary frontier.

Now, while we can establish the non-convergence of the play, fully characterizing the frontier is analytically challenging because its shape will be determined by the optimal play of the game, and the optimal play of the game, in turn, depends on the shape of the frontier. However, the simple assumption of strict concavity allows us to derive some of the basic features of the optimal decision rule, at least around the initial maximizer for the principal, ω_0 :

Proposition 5 *Optimal initial decision rule:* *Around the principal's preferred state, ω_0 , mediocre projects will be favored ($x_{L,\omega_0} > 0$), high-quality projects will be either favored or discriminated against ($x_{H,\omega_0} \gtrless 0$), the admission of mediocre projects is rewarded ($\Delta U_A(A|L, \omega_0) = \Delta U_A(A|R, \omega_0) = \Delta U_A(L, \omega_0) > 0$), as is the acceptance of high-quality alternatives ($\Delta U_A(L, \omega_0) > \Delta U_A(A|H, \omega_0) > 0$), while the rejection of high-quality alternatives is punished ($\Delta U_A(R|H, \omega_0) < 0$).*

Proof. See Appendix A.2 ■

The intuition behind these choices has already been discussed and need not be repeated, with the logic for the distortions in the current state following the logic of the stationary equilibrium while the intuition for the distortions following the admission of mediocre projects and rejection of high-quality projects discussed in conjunction with the agent's truth-telling constraint. The only notable additional result is that the initial acceptance of high-quality projects is also rewarded with future influence. The intuition behind this result is that increasing the influence through $\Delta U_A(A|H, \omega_0)$ is simply another means of deferred compensation. Instead of providing the agent immediate value through distortions in the decision rule to keep him honest, the principal simply promises a higher payoff in the future, which allows her to lower the distortions in the immediate decision rule. Importantly, this is valuable because since the principal can compensate the agent only through distortions in the decision rule, increasingly high payments through the decision rule are increasingly costly due to the high distortions. Thus, instead of providing additional future value simply through a higher $\Delta U_A(L, \omega_0)$, he can claw back some that value by increasing $\Delta U_A(A|H, \omega_0)$ at a lower cost to herself. But because the benefit of $\Delta U_A(L, \omega_0)$ is that it also directly relaxes the truth-telling constraint, the reward offered through $\Delta U_A(A|H, \omega_0)$ is always strictly less.

For the rest of the dynamics, the complex interplay of the constraints makes further characterization more challenging, but the partial results obtained suggests the following conjecture (for a partial discussion, see Appendix A.2):

Conjecture 6 *The behavior of the decision rule away from the principal's preferred state:*

- (i) Both $x_{L,\omega}$ and $x_{H,\omega}$ will be increasing in the current level of influence, ω .
- (ii) $\Delta U_A(L, \omega_0) \geq 0$ and $\Delta U_A(R|H, \omega_0) \leq 0$ will both be decreasing in ω .
- (iii) $\Delta U_A(A|H, \omega_0) \gtrless 0$ will be decreasing in ω and will be negative for high-enough ω .

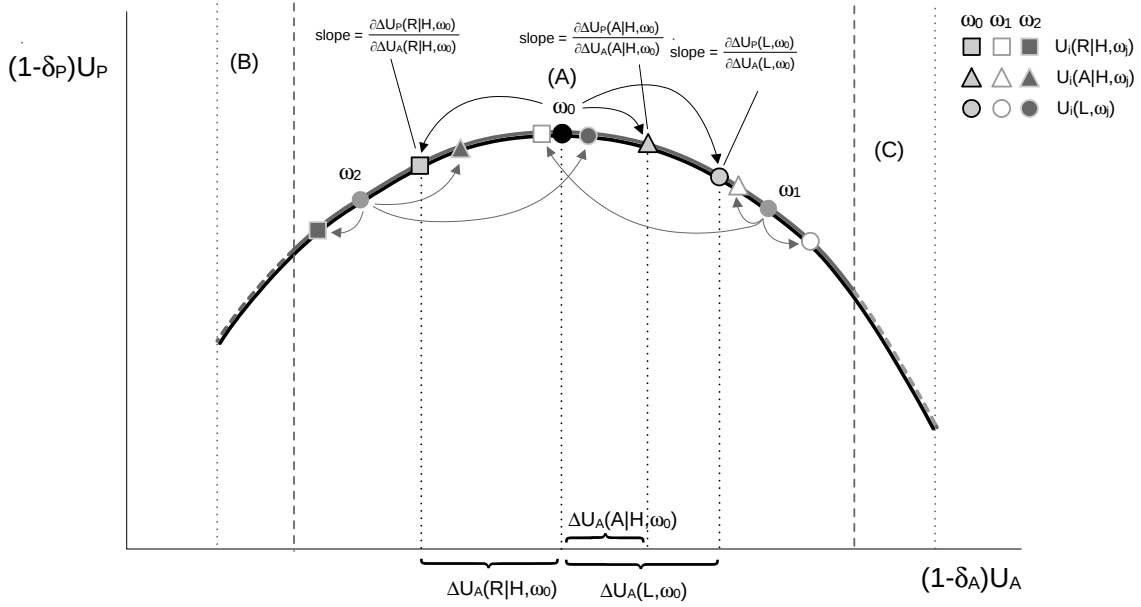


Figure 7: Illustration of the equilibrium dynamics

To build some intuition for this conjecture, consider Figure 7, which illustrates the transitions suggested by the analysis for a potential non-stationary payoff frontier, where segment (A) is the region where only the agent's constraints are binding, while regions (B) and (C) involve a binding constraint by the principal. First, the optimal choice of $x_{L,\omega}$ and $\Delta U_A(L, \omega)$ solves

$$\frac{x_{L,\omega}}{\theta_L} = -\frac{\partial \Delta U_P(L, \omega)}{\partial \Delta U_A(L, \omega)}. \quad (16)$$

Intuitively, the principal can reward the admission of mediocre projects immediately through $x_{L,\omega}$ or delayed reward through continuation value. The current distortion costs the principal $x_{L,\omega}$ while rewarding the agent at rate θ_L , while a future utility costs the principal $\frac{\partial \Delta U_P(L, \omega)}{\partial \Delta U_A(L, \omega)}$. At the margin, the two must be equated. As we increase the agent's influence, moving through states from ω_2 to ω_0 to ω_1 , the more costly the promises of additional future influence become (indeed, at ω_2 , initial additional promise of future influence actually creates value to the principal since we are below the principal's optimum), and thus the additional rewards shrink in size while the immediate settling up through $x_{L,\omega}$ increases. Intuitively, the more influence the principal has already promised to the agent, promising additional influence becomes increasingly costly because it requires increasingly large distortions in the future to honor that promise. Thus, it becomes more cost-effective to begin more aggressive immediate settling up through favorable decisions today.

For $x_{H,\omega}$, the condition is not quite as simple, but the logic is similar. In low states, the distortion is used to penalize the agent while in high states it is used to reward the agent. In addition, note that while the cost of additional rewards through $\Delta U_A(L, \omega_0)$ is increasing in ω , the cost of penalties through $\Delta U_A(R|H, \omega_0)$ is decreasing in ω . Then, for states below the principal's

maximum, the more costly additional punishment through $\Delta U_A(R|H, \omega_0)$ becomes and thus the more the agent is punished immediately through $x_{H, \omega}$. Conversely, for states above the principal's maximum, satisfying the truth-telling constraint through punishment is increasingly attractive as that improves the principal's payoff. Coupling this with the observation that providing additional rewards through $\Delta U_A(L, \omega_0)$ is increasingly costly, the optimal strategy to satisfy the promise-keeping constraint is to use increasingly high immediate settling up through high $x_{H, \omega}$.

Finally, the optimal response in the case of acceptance following the recommendation of a high-quality project is given by

$$\frac{\partial \Delta U_P(A|H, \omega)}{\partial \Delta U_A(A|H, \omega)} = (1 - p) \frac{\partial \Delta U_P(L, \omega)}{\partial \Delta U_A(L, \omega)} + p \frac{\partial \Delta U_P(R|H, \omega)}{\partial \Delta U_A(R|H, \omega)}, \quad (17)$$

so that it is simply a weighted average of the marginal costs of transitions following the admission of a low-quality project and the rejection of a high-quality project. To repeat the logic behind this condition, recall that while this transition has no impact on the truth-telling constraint of the agent, it can be used as an additional means of deferred compensation, benefiting the principal by increasing the continuation value of the game without needing to distort the current decisions. Further, because the expected continuation value is also impacted by the other two transitions, the usefulness of $\Delta U_A(A|H, \omega)$ is determined by its impact on the marginal cost of the other two distortions, $\Delta U_A(L, \omega)$ and $\Delta U_A(R|H, \omega)$.

When the agent's influence is sufficiently low (including the principal's preferred state), this expression implies a positive increase in future influence, suggesting that an agent with lower influence can build up his position by having high-quality projects accepted for implementation. For high levels of influence, however, where $\Delta U_A(L, \omega)$ is small so that $\frac{\partial \Delta U_P(L, \omega)}{\partial \Delta U_A(L, \omega)}$ is close to $\frac{\partial \Delta U_P(\omega)}{\partial \Delta U_A(\omega)}$, then the above condition implies that $\frac{\partial \Delta U_P(A|H, \omega)}{\partial \Delta U_A(A|H, \omega)} > \frac{\partial \Delta U_P(\omega)}{\partial \Delta U_A(\omega)}$. In other words, further acceptances of high-quality projects will actually lower the agent's influence. Intuitively, when the agent is already highly influential, the continuation value is already high and it is relatively easy to sustain truth-telling and additional promises are increasingly costly. As a result, it will be efficient for the principal to have the agent to consume some of that influence and restore the game towards her preferred equilibrium at a limited cost in terms of the additional distortions in $(x_{L, \omega}, x_{H, \omega})$ required to sustain the promise-keeping constraint.

Regions (B) and (C): When the principal's IC constraint becomes binding, those constraints will limit the distortions that can be implemented and thus affect the sustainability of the relationship, as in the stationary case. But as with the agent, changes in continuation value can potentially be used to manage the principal's constraints as well. In region (B), the binding constraint will generally be the rejection of high-quality alternatives. In this case, we could relax the principal's acceptance constraint by actually rewarding a rejection of a high-quality alternative. This, of course, goes against the agent's truth-telling constraint, but it appears feasible that we may be able to relax the principal's IC constraint enough so that we can readjust the stage-game distortions so that also the agent's truth-telling constraint continues to be satisfied. In this case, region (B) would represent a truly probationary period where, following one period of strong discrimination against any

proposals, all outcomes lead to higher influence for the agent in the following period.

In region (C), the binding constraint will generally be the acceptance of high-quality alternatives. Then, the use of $\Delta U_A(A|H, \omega) < 0$, that is, the using up of influence, will have an additional benefit of relaxing the principal's IC constraint by rewarding the acceptance of high-quality alternatives with an increase in her continuation payoff. In other words, the principal will be willing to honor the "favor" that is called in because once the favor is honored, the principal's payoff is improved.

Now, while the general features of the equilibrium appear intuitive, a finer characterization of the optimal strategies is challenging because the key determinant for the transitions is the slope of the payoff frontier at any given point, which is only determined as a part of the whole equilibrium. For example, we do not know whether regions (B) and (C) are always reached as a part of the optimal equilibrium. As a result, to complement this discussion, the following section provides a simplified illustration of some of the forces through a three-state example. Work on a more detailed description of the equilibrium strategies is ongoing.

6.2 A three-state example

This section provides a simple three-state illustration, where the game can be in one of three states, "probationary," "status quo" and "favorable," as ranked by the agent's continuation values in each state, and the principal chooses optimally both the stage-game distortions and the transition probabilities between the states. For computational feasibility, the current illustration makes the simplifying assumption that $\Delta U_A(A|H, \omega) = 0$, so that the acceptance of high-quality alternatives does not affect the equilibrium play.

The resulting equilibrium is illustrated in Figure 8. Panel (a) plots the equilibrium distortions. As expected, the distortions towards both mediocre and high-quality projects are ranked by state, with agent's influence increasing in the state. The only exception to this is that the distortions with respect to mediocre projects match each other for the status quo and probationary states. The reason for this result lies in panel (b), which plots the corresponding transition probabilities: returning from the probationary period involves a sufficiently high expected reward to sustain truth-telling that it transitions with a positive probability to the favorable state. As a result, the marginal cost of increasing the reward for the admission of low-quality projects is the same for status quo and probationary states and, as a result, the optimal distortion is the same as well. As the parties become increasingly impatient, the distortions needed grow, until the principal's IC constraint becomes binding, here initially for the high state, which limits the distortions that can be sustained until the relationship becomes unsustainable, as in the static setting. Panel (b) illustrates that the dynamic setting allows for the additional tool of altering the probability of transitions to sustain some of the incentives, so that the total transition probabilities shoot up until the relationship becomes unsustainable.

The transition probabilities also indirectly illustrate the asymmetric optimal transitions between the states ($q_{ij}(k)$ implies transition from i to j following outcome k). Because of the kink in the

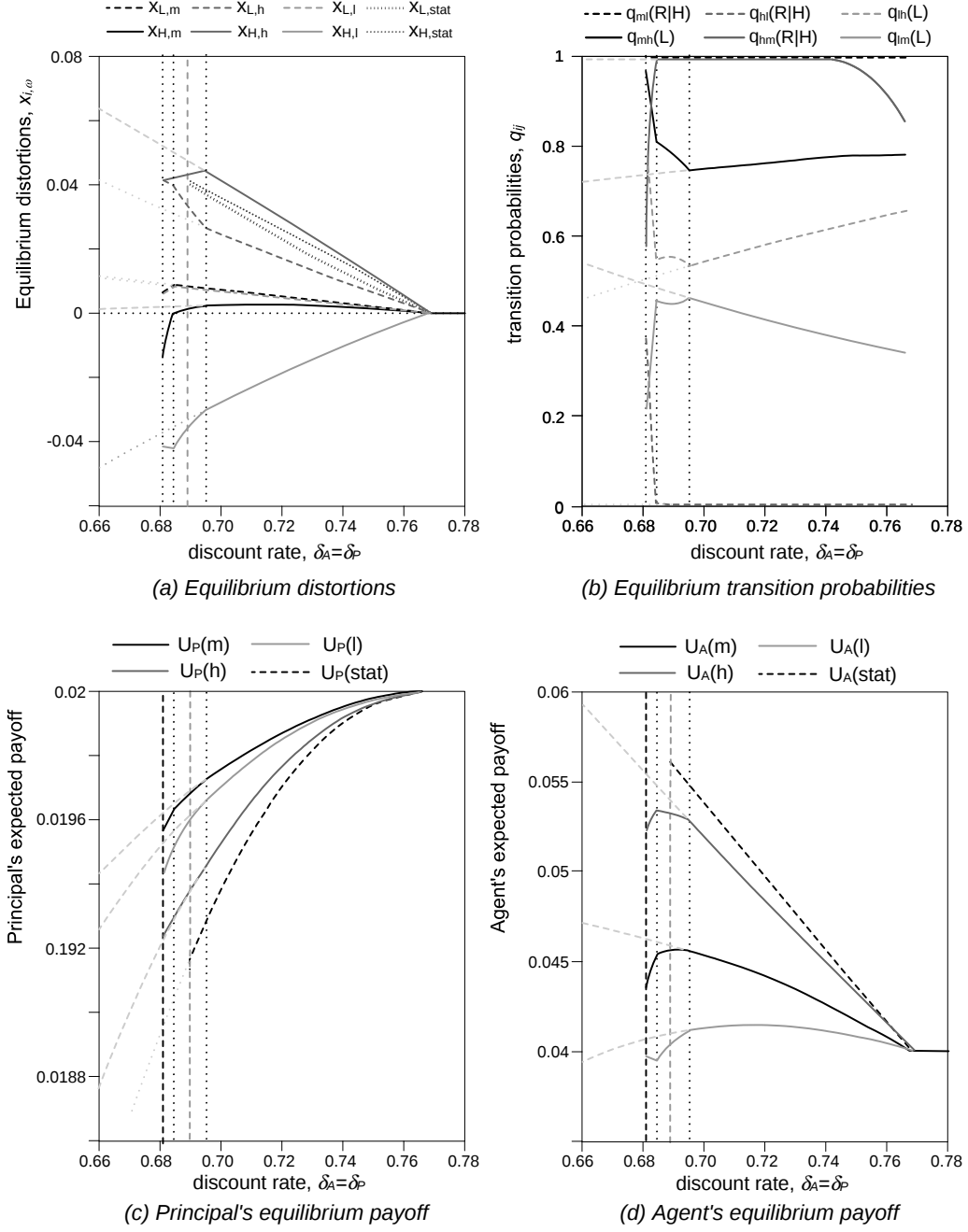


Figure 8: An example of a three-state equilibrium: $\theta_H = 0.6, \theta_L = 0.2$ and $p = 0.5$.

frontier at the status quo state due to the restriction on three states, there is a strictly positive cost for both over- and under-shooting this point. Thus, for example, the game transitions with a probability less than one from the status quo state to the favorable state, so that the game can return to the status quo state with probability 1 instead of needing to overshoot to the probationary state to sustain incentives in the favorable state. For the probationary state, on the other hand, sustaining a sufficient distortion to simply return to the status quo is too costly, so the game switches from the status quo to the probationary state with probability 1, but then overshoots with positive probability and goes directly to the favorable state to optimally sustain truth-telling in that state.

In relation to the stationary equilibrium, panel (a) highlights the benefit of the multi-state strategy by plotting the corresponding stationary equilibrium distortions. The promises and threats through the continuation values allow us to sustain much smaller stage-game distortions especially for the status quo state. This, in turn, sustains the existence of a non-stationary equilibrium for patience levels for which no stationary equilibrium exists. The benefit is further illustrated in panels (c) and (d), which plot the state-dependent continuation values for the principal and the agent. The use of the three-state strategy in this case provides such a high benefit for the principal that her payoff is higher even in her worst state (favorable) than her payoff in the stationary equilibrium. The reason is that, as seen in panel (d), the variation in continuation values allows us to substitute for the level in the continuation value for the agent, so that the agent's payoff is correspondingly lower even in the favorable state than in the stationary equilibrium.

7 Conclusion

I have illustrated how a decision-maker can use the repeated nature of a relationship and the manipulation of her acceptance rule to sustain an informative relationship with an agent. In a stationary equilibrium, to sustain truth-telling, the decision-maker implements a form of corporate socialism, where the decisions are biased in favor of mediocre alternatives to reward honesty. In a non-stationary equilibrium, the principal can further reward the agent through the allocation of future influence, where the admission of mediocre alternatives is rewarded with increased future influence while the rejections of high-quality alternatives are associated with the erosion of influence. In addition, the acceptance of high-quality projects builds up influence for low and intermediate levels of influence, while eroding influence once the agent is sufficiently influential already, reflecting how influence is both build up and used up over time. These results are, however, only a first pass towards a richer understanding of influence in organizations. In terms of dynamics, work remains to be done on what is the truly optimal equilibrium for the principal. In terms of structure, the analysis focused on considering the management of the relationship with a single agent, and an interesting open question is how to manage relationships with multiple agents who provide competing proposals or otherwise multiple relevant pieces of information. Even in terms of variations of the current framework, questions remain, with some final observations made below:

Observability: The analysis assumed a stark asymmetry between perfect observability of the

value of chosen projects and no learning on rejected projects. If there was some probability of learning the value of rejected projects, then the need for a penalty in the case of rejection of high-quality projects would correspondingly decrease. Introducing noisy observability of outcomes for either case would be considerably more challenging, now requiring strategies that depend on the full history of outcomes and comparing the empirical distribution of outcomes with that suggested by the equilibrium strategies.

Timing: The analysis assumed that the agent must make his recommendation before the realization of the principal's state. In terms of timing, this structure is better for the principal than allowing the agent to see the cost as that would tighten the agent's truth-telling constraint by making it clear whether a lie is valuable or not (the key element here is that when choosing which message to send, the agent does not know if misrepresentation is needed to get the project implemented). If the agent could see the principal's state before the recommendation would, however, change the equilibrium dynamics by eliminating the role of rejection of high-quality projects in managing the truth-telling constraint.

Transfers: While deep pockets would eliminate the use of dynamics in managing the relationship, assuming one-sided transfers from the principal to the agent as compensation for truthful reports would leave the basic logic of the results unchanged. In particular, note that ratio of marginal cost and marginal return to money is 1, while compensating the agent through distortions in the decision rule compensate the agent at rate θ_i with marginal cost of x_i . Thus, money would only cap the positive distortions to $x_i \leq \theta_i$, while not helping at all in the case of penalizing the agent.

Randomization: Finally, the analysis focused on truth-telling by the agent followed by deterministic behavior by the principal, and the question remains whether randomization by the agent or the principal could help to improve the outcome. While not immediate, the logic of the framework suggests that randomization will not help. For the principal, the optimal transitions are deterministic, as is the degree of favoritism x_i and the use of a threshold rule, $\underline{c}(m_i)$. Randomization doesn't help to penalize the agent, who only cares about the expected probability of implementation, so the principal should meet these targets in the most efficient manner, which is deterministic. Randomization by the agent, on the other hand, could potentially be used to relax the principal's IC constraint. The most plausible scenario arises when $x_H < 0$ and $|x_H| \geq x_L$ (requiring $p < 1/2$ in the stationary equilibrium). In this case, we could mix some low-quality projects to the high recommendation to relax the constraint, but the downside is that such randomization leads to a loss of information and strictly worse payoff to the principal, potentially lowering the sustainable threshold even more. The principal could, instead, simply use the slack in the other IC constraint to readjust the distortions to restore truth-telling. When $|x_H| = x_L$, the slack disappears and no truthful equilibrium exists. Now, the agent could potentially randomize in both directions to relax the constraint, but again the question remains if the principal's continuation value drops even faster due to such randomization, in particular because now such randomization requires an agent with a high-quality proposal to sometimes claim it is mediocre, which appears highly expensive to the principal (since the constraint is otherwise slack).

References

- [1] Abreu, Dilip, David Pearce and Ennio Stacchetti (1990), "Toward a Theory of Discounted Repeated Games with Imperfect Monitoring," *Econometrica*, 58(5), pp.1041-1063
- [2] Athey, Susan and Kyle Bagwell (2001), "optimal collusion with private information," *rand*, 32(3), 428-65
- [3] Athey, Susan, Kyle Bagwell and Chris Sanchirico (2004), "Collusion and Price Rigidity," *Review of Economic Studies*, 71, 317-349
- [4] Alonso, Ricardo and Niko Matouschek (2007), "Relational Delegation," *The RAND Journal of Economics*, 38(4), 1070-1089
- [5] Andrews, Isaiah and Daniel Barron (2014), "The Allocation of Future Business: Dynamic Relational Contracts with Multiple Agents," working paper, Harvard Society of Fellows and Northwestern Kellogg
- [6] Baker, George, Robert Gibbons and Kevin Murphy (1999), "Informal Authority in Organizations," *Journal of Law, Economics, and Organization*, 15(1), pp. 56-73
- [7] Barron, Daniel (2013), "Putting the Relationship First: Relational Contracts and Market Structure," working paper, Northwestern Kellogg
- [8] Barron, Daniel (2014), "Attaining Efficiency with Imperfect Public Monitoring and One-Sided Markov Adverse Selection," working paper, Northwestern Kellogg
- [9] Barron, Daniel and Michael Powell (2015), "Policies in Relational Contracts," working paper, Northwestern Kellogg
- [10] Board, Simon (2011), "Relational Contracts and the Value of Loyalty," *American Economic Review*, 101, 3349-3367
- [11] Calzolari, Giacomo and Giancarlo Spagnolo (2010), "Relational Contracts and Competitive Screening," CEPR Discussion Paper No. 7434
- [12] Campbell, Arthur (2015), "Political Capital," working paper, Yale University
- [13] Chakraborty, Archishman and Bilge Yilmaz (2013), "Authority, Consensus and Governance," working paper, Yeshiva University and University of Pennsylvania
- [14] Crawford, Vincent and Joel Sobel (1982) "Strategic Information Transmission," *Econometrica*, 50, 1431-1451
- [15] Garfagnini, Umberto, Marco Ottaviani and Peter Sorensen (2014), "Accept or Reject? An Organizational Perspective," *International Journal of Industrial Organization*, 34, pp. 66-74
- [16] Guo, Yingni and Johannes Horner (2015), "Dynamic Mechanisms without Money," Cowles foundation discussion paper No. 1985
- [17] Hauser, Christine and Hugo Hopenhayn (2005), "Trading Favors: Optimal Exchange and Forgiveness," working paper, University of Rochester and UCLA

- [18] Kolotilin, Anton and Hongyi Li (2015), "Relational Communication with Transfers," working paper, University of New South Wales
- [19] Levin, Jonathan (2002), "Multilateral Contracting and the Employment Relationship," *The Quarterly Journal of Economics*, pp. 1075-1103
- [20] Levin, Jonathan (2003), "Relational Incentive Contracts," *American Economic Review*, 93(3), 835-857
- [21] Li, Jin and Niko Matouschek (2013), "Managing Conflicts in Relational Contracts," *AER* 103(6), pp.2328-2351
- [22] Li, Jin, Niko Matouschek and Michael Powell (2015), "Power Dynamics in Organizations," working paper, Northwestern University
- [23] Li, Zhouzheng, Heikki Rantakari and Huanxing Yang (2016) "Competitive Cheap Talk," *Games and Economic Behavior*, 96, pp. 65-89
- [24] Lipnowski, Elliot and Joao Ramos (2015), "Repeated Delegation," working paper, New York University
- [25] Malcomson, James (2015), "Relational Incentive Contracts with Private Information," University of Oxford Department of Economics Discussion Paper Series No. 634
- [26] Mobius, Markus (2001), "Trading Favors," working paper
- [27] Padro i Miquel, Gerard and Pierre Yared (2012), "The Political Economy of Indirect Control," *QJE* 127, 947-1015
- [28] Rantakari, Heikki (2015), "Managerial Influence and Organizational Performance," working paper, University of Rochester
- [29] Rantakari, Heikki (2016), "Soliciting Advice: Active versus Passive Principals," forthcoming, *Journal of Law, Economics, and Organization*
- [30] Renault, Jérôme, Eilon Solan and Nicolas Vieille (2013), "Dynamic sender-receiver games," *Journal of Economic Theory*, 148(2), pp. 502-534

A Proofs and derivations

A.1 Solving for the stationary equilibrium

I will perform the analysis of the stationary solution in two steps. First, I will consider the solution when only the agent's truth-telling constraint is binding. Second, I will introduce the principal's incentive-compatibility constraint for implementing the required distortion.

For the first step, we need to consider the principal's objective function and the agent's truth-telling constraint. The basic features are outlined next.

Principal's objective function: The principal's objective function is to minimize the loss $px_H^2 + (1-p)x_L^2$. Since x_L will be positive in equilibrium, we can characterize the principal's indifference curves defined by the ellipses $x_H = \pm \sqrt{\frac{K-(1-p)x_L^2}{p}}$. For later, note that the slope of the indifference curves is given by $\frac{dx_H}{dx_L} = \pm \left| \frac{(1-p)x_L}{p} \sqrt{\frac{p}{K-(1-p)x_L^2}} \right|$ and which equals $\frac{dx_H}{dx_L} = \pm \left| \frac{1-p}{p} \right|$ at $x_H = x_L = \sqrt{K}$. Relatedly, the second important observation is that the slope of the indifference curve is constant along any ray $x_H = \alpha x_L$, given by

$$\frac{dx_L}{dx_H} = \pm \left| \frac{p}{(1-p)\alpha} \right|. \quad (18)$$

Agent's truth-telling constraint: Letting $y_A = \frac{1-\delta_A}{\delta_A}$, we can rearrange the agent's truth-telling constraint given by equation 8 as

$$x_L = \frac{y_A [(\theta_H + x_H) - (\theta_L)] \theta_L - (\theta_H + x_H) (\phi + p\theta_H x_H)}{\theta_L [y_A + (\theta_H + x_H) (1-p)]}. \quad (19)$$

or

$$x_H = \frac{- (\phi + (1-p) \theta_L x_L - y_A \theta_L + p\theta_H^2) \pm \sqrt{((\phi + (1-p) \theta_L x_L - y_A \theta_L) - p\theta_H^2)^2 - 4p\theta_H \theta_L y_A (\theta_L + x_L)}}{2p\theta_H}. \quad (20)$$

Because increasing x_L both increases continuation value and reduces reneging temptation, x_L will always be non-negative. The distortion for the high-quality project, x_H , on the other hand, may be positive or negative, to be examined in more detail below.

Properties of the solution: The solution involves the principal choosing the (x_L, x_H) that reaches the indifference curve closest to the origin, subject to the agent's truth-telling constraint. As illustrated in Figure 3, this involves the standard tangency condition between the indifference curve and the constraint. As drawn, it leads to a clockwise-rotation in the (x_H, x_L) -space, originating

from $x_L = x_H = 0$ and converging to $x_H = -(1-p)(\theta_H - \theta_L)$, $x_L = p(\theta_H - \theta_L)$. This section will analytically show that the solution will always take the shape as drawn in the Figure.

The preliminary observations are as follows. First, from equation 20, we can establish that (i) as $\delta_A \rightarrow 1$, the constraint converges to $x_H \geq -\frac{(\phi+(1-p)\theta_L x_L)}{p\theta_H}$, (ii) as $\delta_A \rightarrow 0$, the solution converges to $x_H \leq -(\theta_H - \theta_L) + x_L$ and (iii) all solutions (including $\delta_A \rightarrow 1$ and $\delta_A \rightarrow 0$) pass through $x_H = -(1-p)(\theta_H - \theta_L)$, $x_L = p(\theta_H - \theta_L)$. Intuitively, the decision rule where all information is ignored ($\theta_H + x_H = \theta_L + x_L = E(\theta_i)$) is always incentive-compatible to the agent. Second, from equation 19 we get that

$$\frac{\partial x_L}{\partial y_A} = \frac{(\theta_H + x_H)E(\theta_i)((\theta_H + x_H) - E(\theta_i))}{\theta_L((1-p)(\theta_H + x_H) + y_A)^2}, \quad (21)$$

so that increased impatience and given x_H requires an increase in x_L for any $x_H > -(1-p)(\theta_H - \theta_L)$ and a decrease for any $x_H < -(1-p)(\theta_H - \theta_L)$. This implies that the feasible solutions lie either in the cone satisfying $x_H \geq -\frac{(\phi+(1-p)\theta_L x_L)}{p\theta_H}$ and $x_H \geq -(\theta_H - \theta_L) + x_L$, or in the cone satisfying both $x_H \leq -\frac{(\phi+(1-p)\theta_L x_L)}{p\theta_H}$ and $x_H \leq -(\theta_H - \theta_L) + x_L$. Third, the principal's indifference curve has a slope of 1 at $(x_H = -(1-p)(\theta_H - \theta_L), x_L = p(\theta_H - \theta_L))$, which matches the slope of the truth-telling constraint for $\delta_A \rightarrow 0$. As a result, the principal can never be worse off than ignoring the information provided, and so we can rule out the second region. Thus, the solution will always lie in the region of $x_H \geq -(1-p)(\theta_H - \theta_L)$. Fourth, note further that

$$\frac{\partial^2 x_L}{\partial y_A \partial x_H} = \frac{E(\theta_i)(E(\theta_i)(1-p)[\theta_H + x_H] + y_A(2(\theta_H + x_H) - E(\theta_i)))}{\theta_L((1-p)(\theta_H + x_H) + y_A)^3} \geq 0, \quad (22)$$

so that in the relevant region, the effect of the agent's discount rate is increasing in the level of favoritism towards the agent with high-quality proposals. This creates the "fanning out" of the truth-telling constraints as illustrated in Figure 3. The basic shape of the truth-telling constraints thus matches the Figure.

To complete the characterization, the logic behind the proof is illustrated in Figure 9. First, recall that the slope of the principal's indifference curve is constant along any ray from the origin, and the optimal solution is characterized by the tangency point of the truth-telling constraint with the indifference curve. Then, if we can establish that the slope of the truth-telling constraint along the ray satisfies $\frac{d}{dy_A} \left(\frac{dx_L}{dx_H} \Big|_{x_H=\alpha x_L} \right) > 0$, then there is exactly one truth-telling constraint that has a slope matching that of the indifference curves along that ray, and any lower truth-telling constraint (more patient agent) will have a solution on the ray counter-clockwise while any higher truth-telling constraint (less patient agent) will have the solution on a ray that is clockwise from the current ray. In the Figure itself, when the the agent is sufficiently patient, the solution lies on ray A, given by point 3. Given the rotation $\frac{d}{dy_A} \left(\frac{dx_L}{dx_H} \Big|_{x_H=\alpha x_L} \right) > 0$, the truth-telling constraints for a less patient agent (points 1 and 2) intersect the corresponding indifference curve from below, and the same for all rays counter-clockwise from the ray, so that no tangency point exists. As a result, the new solution must lie clockwise from ray A for any less patient agent. When the agent is somewhat less patient

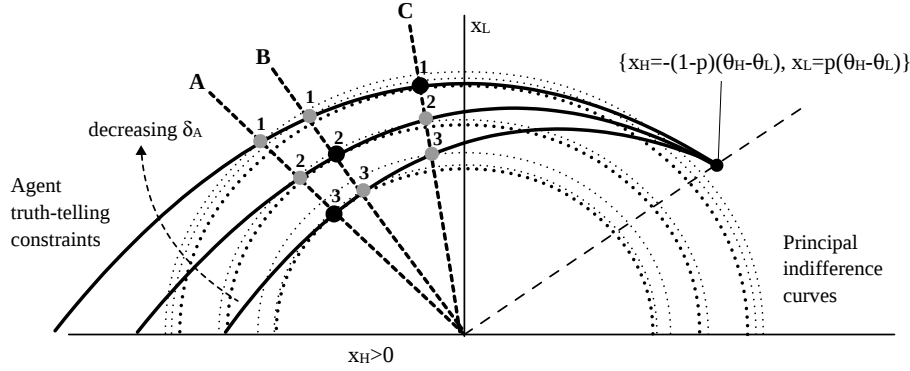


Figure 9: Intuition behind the solution

(ray B), the tangency is found at point 2, with the rotation implying that the truth-telling constraint of the less patient agent intersects the indifference curve from below (point 1), while a more patient agent intersects from above (point 3). Finally, for the least patient agent, the tangency condition is found along ray C (point 1), while the truth-telling constraints for the relatively more patient agents intersect from above, implying their tangency points must lie counter-clockwise from the ray. The remaining task is simply to establish $\frac{d}{dy_A} \left(\frac{dx_L}{dx_H} \Big|_{x_H=\alpha x_L} \right)$, where the challenge is that we are not only changing the slope of the truth-telling constraint but also the location of the truth-telling constraint.

As a preliminary step, consider the solution when the truth-telling constraint becomes just binding at the first-best solution of $x_L = x_H = 0$, which arises when $y_A \rightarrow \underline{y}_A = \frac{p(1-p)(\theta_H - \theta_L)\theta_H}{\theta_L}$. From equation 20 we can derive that the slope of the truth-telling constraint can be written as

$$\frac{dx_H}{dx_L} = - \frac{\theta_L ((1-p)(\theta_H + x_H) + y_A)}{(2p\theta_H x_H + (\phi + (1-p)\theta_L x_L - y_A\theta_L + p\theta_H^2))}, \quad (23)$$

which simplifies, for $x_L = x_H = 0$ and $\underline{y}_A = \frac{p(1-p)(\theta_H - \theta_L)\theta_H}{\theta_L}$ to

$$\frac{dx_H}{dx_L} = - \frac{(1-p)\theta_H E(\theta)}{p\theta_H E(\theta) + \phi(\theta)}. \quad (24)$$

Then, using the tangency condition with the principal's indifference curve we can identify the first ray characterizing the solution through $-\frac{(1-p)}{p\alpha} = -\frac{(1-p)\theta_H E(\theta)}{p\theta_H E(\theta) + \phi}$, giving

$$\lim_{y_A \rightarrow \underline{y}_A} \left(\frac{x_H}{x_L} \right) = 1 + \frac{(1-p)(\theta_H - \theta_L)^2}{\theta_H E(\theta)} < \frac{\theta_H}{\theta_L}. \quad (25)$$

This expression identifies the slope of the equilibrium (x_H, x_L) path at the origin. Then, to establish how the slope along the rays changes, we can write $\frac{d}{dy_A} \left(\frac{dx_L}{dx_H} \Big|_{x_H=\alpha x_L} \right)$ as

$$\frac{\partial^2 x_L}{\partial x_H \partial y_A} + \frac{\partial^2 x_L}{\partial^2 x_H} \frac{dx_H}{dx_L} \frac{dx_L}{dy_A}, \quad (26)$$

where the first component identifies the rotation of the indifference curve, while the second component identifies the change in the point around which the slope is evaluated. From above we already know that

$$\frac{\partial^2 x_L}{\partial y_A \partial x_H} = \frac{E(\theta)(E(\theta)(1-p)[\theta_H + x_H] + y_A(2(\theta_H + x_H) - E(\theta)))}{\theta_L((1-p)(\theta_H + x_H) + y_A)^3}, \quad (27)$$

while given the assumption of the ray, $\frac{dx_H}{dx_L} = \alpha$. Next, we can use equation 19 to relate x_L and y_A along the ray as

$$x_L = \frac{-((\alpha-1)(p\theta_H^2 - \theta_L y_A) + \theta_H E(\theta) + \alpha\phi) + \sqrt{((\alpha-1)(p\theta_H^2 - \theta_L y_A) + \theta_H E(\theta) + \alpha\phi)^2 + 4\alpha(\theta_L y_A(\theta_H - \theta_L) - \phi\theta_H)((\alpha-1)p\theta_H + E(\theta))}}{2\alpha((\alpha-1)p\theta_H + E(\theta))},$$

which gives us the slope $\frac{dx_L}{dy_A}$ as

$$\frac{dx_L}{dy_A} = \frac{\theta_L((\alpha-1)x_L + (\theta_H - \theta_L))}{2\alpha((\alpha-1)p\theta_H + E(\theta_i))x_L + ((\alpha-1)(p\theta_H^2 - \theta_L y_A) + \theta_H E(\theta_i) + \alpha\phi)} \geq 0.^8 \quad (28)$$

Finally, to get $\frac{\partial^2 x_L}{\partial^2 x_H}$, we can use equation 19 to get

$$\frac{\partial^2 x_L}{\partial^2 x_H} = -\frac{2y_A E(\theta)((1-p)E(\theta) + y_A)}{\theta_L((1-p)(\theta_H + x_H) + y_A)^3} < 0. \quad (29)$$

Then, bringing the components together, we have

$$\frac{d}{dy_A} \left(\frac{dx_L}{dx_H} \Big|_{x_H=\alpha x_L} \right) = \underbrace{\frac{\partial^2 x_L}{\partial x_H \partial y_A}}_{>0} + \underbrace{\frac{\partial^2 x_L}{\partial^2 x_H}}_{<0} \underbrace{\frac{dx_H}{dx_L}}_{\geq 0} \underbrace{\frac{dx_L}{dy_A}}_{>0}. \quad (30)$$

When $\frac{dx_H}{dx_L} < 0$ ($\alpha \leq 0$), the solution is thus immediate because the shift in location reinforces the original rotation through $\frac{\partial^2 x_L}{\partial x_H \partial y_A}$. For $\alpha > 0$, we need to complete the expression. Substituting in the components, we get

$$\frac{E(\theta)}{\theta_L((1-p)(\theta_H + x_H) + y_A)^3} \left[\frac{(E(\theta)(1-p)[\theta_H + x_H] + y_A(2(\theta_H + x_H) - E(\theta)))}{1} - \frac{2\alpha y_A \theta_L((x_H - x_L) + (\theta_H - \theta_L))((1-p)E(\theta) + y_A)}{\eta(2x_H + \theta_H) + \alpha\phi - (\alpha-1)\theta_L y_A} \right],$$

where $\eta = (\alpha-1)p\theta_H + E(\theta)$. Now, substituting $y_A = \frac{(\theta_H + x_H)u_A}{((x_H - x_L) + (\theta_H - \theta_L))\theta_L}$ for the first component and $y_A \theta_L((x_H - x_L) + (\theta_H - \theta_L)) = (\theta_H + x_H)u_A$ for the second component by using the agent's truth-telling constraint, the expression simplifies further, after some algebra, to

$$\frac{E(\theta)(\theta_H + x_H)}{\theta_L((1-p)(\theta_H + x_H) + y_A)^3} \left[\frac{((\theta_H - \theta_L)(1-p) + x_H)[E(\theta)^2 + u_A]}{((x_H - x_L) + (\theta_H - \theta_L))\theta_L} - \frac{2\alpha((1-p)E(\theta) + y_A)u_A}{\eta(2x_H + \theta_H) + \alpha\phi - (\alpha-1)\theta_L y_A} \right].$$

Next, substituting $\alpha = \frac{x_H}{x_L}$ gives the expression as

⁸While not immediately obvious from the expression, this result follows from equation 21, implying that the truth-telling constraints do not intersect in the relevant region.

$$\frac{E(\theta)(\theta_H+x_H)}{\theta_L((1-p)(\theta_H+x_H)+y_A)^3} \left[\frac{((\theta_H-\theta_L)(1-p)+x_H)E(\theta)^2}{((x_H-x_L)+(\theta_H-\theta_L))\theta_L} + 2u_A \left(-\frac{\frac{((\theta_H-\theta_L)(1-p)+x_H)}{((x_H-x_L)+(\theta_H-\theta_L))\theta_L}}{\frac{x_H((1-p)E(\theta)+y_A)}{[px_H\theta_H+(1-p)x_L\theta_L](x_H+\theta_H)+x_Hu_A-(x_H-x_L)\theta_Ly_A}}} \right) \right].$$

Now, the first two components are positive, so the derivative is positive as long as

$$\frac{((\theta_H-\theta_L)(1-p)+x_H)}{((x_H-x_L)+(\theta_H-\theta_L))\theta_L} - \frac{x_H((1-p)E(\theta)+y_A)}{[px_H\theta_H+(1-p)x_L\theta_L](x_H+\theta_H)+x_Hu_A-(x_H-x_L)\theta_Ly_A} \geq 0.$$

The final step is to substitute $y_A = \frac{(\theta_H+x_H)u_A}{((x_H-x_L)+(\theta_H-\theta_L))\theta_L}$ for the second component, after which the expression, after some tedious algebra, finally simplifies to

$$\frac{E(\theta)((\theta_H-\theta_L)(1-p)+x_H)^2(\theta_Hx_L-\theta_Lx_H)}{\theta_L((x_H-x_L)+(\theta_H-\theta_L))[(\theta_H+x_H)(\theta_H-\theta_L)E(\theta)E(x)+x_H[(\theta_H-\theta_L)+(x_H-x_L)]u_A]} \geq 0,$$

where the result follows from the observation that we only need to consider the region of $x_H > 0$ ($x_H \leq 0$ was already established earlier), $E(x) \geq 0$ due to the bounds on the potentially optimal solutions, and $\frac{\theta_H}{\theta_L} > \frac{x_H}{x_L}$ from the preliminary result of the optimal solution when the truth-telling constraint becomes binding. This completes the proof.

Two important corollaries follow. First, the clockwise rotation implies that $\frac{d(\frac{x_H}{x_L})}{dy_A} < 0$, so that the ratio of favoritism between high- and low-quality projects is monotone decreasing in the impatience of the agent. Second, we can write the agent's payoff as

$$\phi + p\theta_Hx_H + (1-p)\theta_Lx_L = \phi + x_L \left(p\theta_H \left(\frac{x_H}{x_L} \right) + (1-p)\theta_L \right),$$

so that we can write the change in the agent's payoff as

$$\frac{dx_L}{dy_A} \left(p\theta_H \left(\frac{x_H}{x_L} \right) + (1-p)\theta_L \right) + x_L \left(p\theta_H \frac{d(\frac{x_H}{x_L})}{dy_A} \right),$$

which is positive for a sufficiently patient agent ($\frac{dx_L}{dy_A} > 0$, $\frac{x_H}{x_L} > 0$ and x_L is small), while negative for a sufficiently impatient agent ($\frac{dx_L}{dy_A} < 0$).

Adding the principal's IC constraint: Having established the solution in the absence of the principal's constraint, we can now add the principal's IC constraint, given by

$$|x_i| \leq \frac{\delta_P}{2(1-\delta_P)} (\phi - px_H^2 - (1-p)x_L^2).$$

Letting $y_P = \frac{(1-\delta_P)}{\delta_P}$, we can write the maximum distortions as

$$\begin{aligned} \bar{x}_L &= \frac{\sqrt{y_P^2 + (1-p)(\phi - px_H^2)} - y_P}{(1-p)} & \text{for } x_L > |x_H| \\ \bar{x}_H &= \pm \frac{\sqrt{y_P^2 + p(\phi - (1-p)x_L^2)} - y_P}{p} & \text{for } |x_H| > x_L \end{aligned}$$

For later, an important benchmark is given by $\bar{x}_L = \bar{x}_H$, which gives the maximum sustainable distortion as $\bar{x}_L = \bar{x}_H = \sqrt{y_P^2 + \phi} - y_P$. Now, index the indifference curves so that the IC constraint of the principal is paired with an indifference curve that touches the constraint at $\bar{x}_L = \bar{x}_H$. Then,

we have that the indifference curve $x_H = \pm \sqrt{\frac{K-(1-p)x_L^2}{p}}$ that matches the IC constraint is given by $K = \left(\sqrt{y_P^2 + \phi} - y_P\right)^2$. Then, we obtain that

$$\left| \frac{dx_H}{dx_L} \Big|_{|x_H| > x_L} \right| = \frac{(1-p)x_L}{\sqrt{y_P^2 + p(\phi - (1-p)x_L^2)}} < \frac{(1-p)x_L}{p} \sqrt{\frac{p}{K - (1-p)x_L^2}}$$

and

$$\left| \frac{dx_L}{dx_H} \Big|_{|x_H| < x_L} \right| = \frac{px_H}{\sqrt{y_P^2 + (1-p)(\phi - px_H^2)}} < \left| \frac{dx_L}{dx_H} \right| = \frac{px_H}{(1-p)\sqrt{\frac{K - px_H^2}{1-p}}},$$

which implies that the slope of the IC constraint is always flatter than the corresponding portion of the indifference curve. When the IC constraint touches the indifference curve ($|x_H| \rightarrow x_L$), we have

$$\left| \frac{dx_H}{dx_L} \Big|_{|x_H| > x_L} \right| > \frac{1-p}{p} > \left| \frac{dx_H}{dx_L} \Big|_{|x_H| < x_L} \right|,$$

so that the IC constraint is kinked at the point where touching the indifference curve. The conclusion is then that the IC constraints can be enclosed inside a given indifference curve, where, for $\delta_P \rightarrow 1$, the slopes coincide and the IC constraint becomes the indifference curve, while for $\delta_P \rightarrow 0$, the slopes of the IC constraint become zero and the IC constraint becomes a square (concentrated around zero).

To verify the exact effects of the constraints, we resort to graphic proof, as illustrated in Figure 10. Panel (i) illustrates the logic when the constraint becomes binding in the region of $x_L > x_H > 0$. When the agent is sufficiently patient (A), the principal's IC constraint is irrelevant. When the agent becomes more impatient, there is a point at which the principal's IC constraint becomes just binding at the optimal solution (B). After this, further impatience of the agent would expand the optimal solution to (C'), but that lies outside the principal's IC constraint, and is thus infeasible. Instead, the solution travels along the principal's IC constraint to (C). If the agent is any more impatient, no solution exists. This solution follows because we know that when the principal's IC becomes binding, the principal's IC constraint intersects the tangency point from below, so that the feasible region remains positive size, until the agent becomes sufficiently impatient. Note that the same logic applies when $x_H < 0$ and $|x_H| > x_L$ simply by rotating the figure 90 degrees clockwise. The only caveat is that this requires that $p < 1/2$, so that the pivot point $x_H = -(1-p)(\theta_H - \theta_L)$, $x_L = p(\theta_H - \theta_L)$ lies in the region. Panel (ii) illustrates the logic when the constraint becomes binding when $x_H > x_L > 0$. Following the same logic, at (B) the constraint becomes just binding, and now intersects the tangency point from above due to its shape, and so the feasible set remains positive and a distorted solution exists. As before, however, the optimal path would take us to (C'), while the requirement to satisfy the principal's IC constraint leads us to (C), after which no further solutions exist. And again, by rotating the picture by 90 degrees clockwise, we get the same logic for $x_H < 0$ and $|x_H| < x_L$, as long as $p < 1/2$ (with the caveat that the set of feasible solutions may disappear before the diagonal is reached due to the relative slopes of the truth-telling and IC constraints).

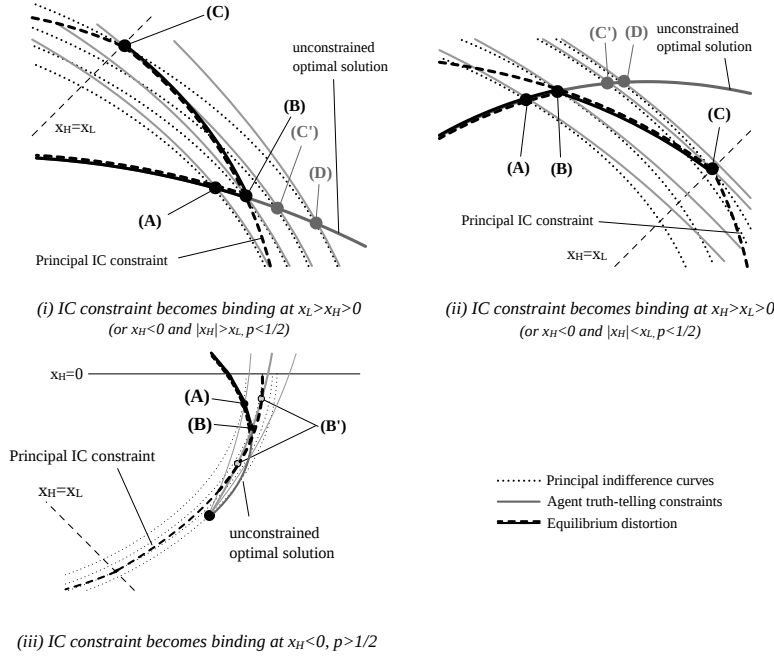


Figure 10: Effects of the principal's IC constraint

In contrast, when $p > 1/2$, no solution exists once the IC constraint becomes binding (note that in this case, $x_H < 0$ and $|x_H| > x_L$ can never arise). This result is illustrated in panel (iii). First, recall that the truth-telling constraints originate from the pivot point, which lies on the outermost indifference curve, and have a positive slope for the region for which the optimal x_H may be negative. Second, recall that any IC constraint can be associated with an indifference curve where it touches it at $x_H = x_L$. Further, the IC constraint is orthogonal to the axis when it intersects it. This means that (a) the IC constraint begins on the diagonal from a lower indifference curve than the truth-telling constraint and (b) it is steeper than the truth-telling constraint once reaching the horizontal axis. This means that the IC constraint always intersects the truth-telling constraint twice. For example, at solution (B), which is feasible, the IC constraint intersects the truth-telling constraint at points (B'). As the agent becomes increasingly impatient, the last tangency point occurs at the last point that the IC and the truth-telling constraint are just satisfied, after which no feasible solution exists.

A.2 Analysis of the non-stationary framework

This subsection contains observations on the behavior of the optimal non-stationary strategies, while recognizing that the section is still very much work in progress. Combining the agent's truth-telling and incentive-compatibility constraints, allows us to write the level of current distortions needed as function of the continuation values as

$$\Pr(A|H, \omega) = \frac{U_A(\omega) + E(\theta)^2 - \delta_A U_A(R|H, \omega)}{E(\theta) + \delta_A p U_A(A|H, \omega) - \delta_A U_A(R|H, \omega)} \quad (31)$$

and

$$\Pr(A|L, \omega) = \frac{U_A(\omega) + E(\theta)^2 - \delta_A U_A(L, \omega) - \Pr(A|H, \omega)(p(\theta_H - \theta_L) + \delta_A p U_A(A|H, \omega))}{\theta_L}, \quad (32)$$

or, equivalently, we can write the needed changes in the continuation values following the admission of a mediocre project and the rejection of a high-quality project as

$$\Delta U_A(R|H, \omega) = \frac{U_A(\omega)(1 - \delta_A) + E(\theta)^2 - \Pr(A|H, \omega)(E(\theta) + \delta_A p \Delta U_A(A|H, \omega) - \delta_A(1 - p)U_A(\omega))}{\delta_A \Pr(R|H, \omega)} \quad (33)$$

and

$$\Delta U_A(L, \omega) = \frac{U_A(\omega)(1 - \delta_A) + E(\theta)^2 - \Pr(A|L, \omega)\theta_L - \Pr(A|H, \omega)(p(\theta_H - \theta_L) + \delta_A p(U_A(\omega) + \Delta U_A(A|H, \omega)))}{\delta_A}. \quad (34)$$

From these expressions we get the following relationships, to be used more below:

$$\begin{aligned} \frac{\partial \Delta U_A(R|H, \omega)}{\partial \Delta U_A(A|H, \omega)} &= -\frac{p \Pr(A|H, \omega)}{\Pr(R|H, \omega)} < 0 \\ \frac{\partial \Delta U_A(R|H, \omega)}{\partial x_L} &= 0 \\ \frac{\partial \Delta U_A(R|H, \omega)}{\partial x_H} &= \frac{E(\theta)^2 - E(\theta) + U_A(\omega) - \delta_A p(U_A(\omega) + \Delta U_A(A|H, \omega))}{\delta_A \Pr(R|H, \omega)^2} < 0 \\ \frac{\partial \Delta U_A(R|H, \omega)}{\partial U_A(\omega)} &= \frac{[1 - \delta_A(1 - (1 - p) \Pr(A|H, \omega))]}{\delta_A \Pr(R|H, \omega)} > 0 \\ \frac{\partial \Delta U_A(L, \omega)}{\partial \Delta U_A(A|H, \omega)} &= -p \Pr(A|H, \omega) < 0 \\ \frac{\partial \Delta U_A(L, \omega)}{\partial x_L} &= -\frac{\theta_L}{\delta_A} < 0 \\ \frac{\partial \Delta U_A(L, \omega)}{\partial x_H} &= -\frac{p(\delta_A(U_A(\omega) + \Delta U_A(A|H, \omega)) + (\theta_H - \theta_L))}{\delta_A} < 0 \\ \frac{\partial \Delta U_A(L, \omega)}{\partial U_A(\omega)} &= \frac{[1 - \delta_A(1 + p \Pr(A|H, \omega))]}{\delta_A} \geq 0^9 \end{aligned}$$

In other words, for the penalty following a rejection of a high-quality project, the penalty is increasing in the reward following the acceptance of a high-quality project, unaffected by the treatment of low-quality projects, increasing in the current level of favoritism towards high-quality projects and decreasing in the promised utility. Intuitively, all follow from the need to satisfy the two constraints: increased future or current reward increases the current value and thus requires a corresponding penalty, low-quality project treatment affects the truth-telling constraint and the continuation value in a way that leaves the future penalty needed unchanged, and a higher promised utility can be reached only by having a lower expected punishment. For the future reward for admitting a low-quality project, the same logic carries through: increased reward through either x_L or $\Delta U_A(A|H, \omega)$

⁹Around the equilibrium

reduces the future reward to maintain the promise-keeping constraint, and the same for increased reward through x_H . Conversely, a higher promised utility increases the need to reward the admittance of low-quality projects.

All the signs are immediate, except for $\frac{\partial \Delta U_A(R|H, \omega)}{\partial x_H}$. But here, note that for the condition to hold, we must have

$$E(\theta)^2 - E(\theta) + U_A(\omega) - \delta_A p (U_A(\omega) + \Delta U_A(A|H, \omega)) < 0.$$

But from $\Delta U_A(R|H, \omega) \leq 0$ we obtain that

$$[(U_A(\omega) \delta_A - E(\theta)) - \delta_A p (U_A(\omega) + \Delta U_A(A|H, \omega))] < -\frac{(1-\delta_A)(V_A(\omega))}{\Pr(A|H, \omega)},$$

where $(1-\delta_A)V_A(\omega) = (1-\delta_A)U_A(\omega) + E(\theta)^2 = (1-\delta_A)(U_A(\omega) + V^{dev})$. Substituting back in the derivative, we get

$$-\frac{(1-\delta_A)(V_A(\omega))}{\Pr(A|H, \omega)} + V_A(\omega)(1-\delta_A) < 0 \Leftrightarrow 0 < \frac{\Pr(R|H, \omega)}{\Pr(A|H, \omega)},$$

which is true. Below is a collection of observations on the properties of the solution.

Lemma 7 *When the principal's IC is slack, then $\Delta U_P(A|L, \omega) = \Delta U_P(R|L, \omega)$:*

For the choice of $\Delta U_P(A|L, \omega)$ and $\Delta U_P(R|L, \omega)$, the principal wants to minimize

$$\delta_P(1-p)(\Pr(A|L, \omega)\Delta U_P(A|L, \omega) + \Pr(R|L, \omega)\Delta U_P(R|L, \omega))$$

subject to the agent receiving a constant continuation value. Optimality for the principal requires that

$$\delta_P(1-p) \left(\Pr(A|L, \omega) \frac{d\Delta U_P(A|L, \omega)}{d\Delta U_A(A|L, \omega)} + \Pr(R|L, \omega) \frac{d\Delta U_P(R|L, \omega)}{d\Delta U_A(R|L, \omega)} \frac{\Delta U_A(R|L, \omega)}{\Delta U_A(A|L, \omega)} \right) = 0,$$

while the agent's constraint implies that $\frac{d\Delta U_A(R|L, \omega)}{d\Delta U_A(A|L, \omega)} = -\frac{\Pr(A|L, \omega)}{\Pr(R|L, \omega)}$,

which simplifies the principal's constraint to

$$\delta_P(1-p) \left(\frac{d\Delta U_P(A|L, \omega)}{d\Delta U_A(A|L, \omega)} - \frac{d\Delta U_P(R|L, \omega)}{d\Delta U_A(R|L, \omega)} \right) = 0,$$

which is satisfied only when $\frac{d\Delta U_P(A|L, \omega)}{d\Delta U_A(A|L, \omega)} = \frac{d\Delta U_P(R|L, \omega)}{d\Delta U_A(R|L, \omega)}$ which, due to the concavity of the frontier, is satisfied only when $\Delta U_A(A|L, \omega) = \Delta U_A(R|L, \omega)$.

Lemma 8 *If the optimal non-stationary equilibrium does not reach the frontier of the full payoff set, then there is no absorbing state*

The proof for this lemma follows from a number of separate observations, as we need to verify that a movement away from the stationary equilibrium is beneficial under different configurations of binding constraints. For intermediate states, we need to worry about the agent's truth-telling constraint, while for the extreme states we need to worry also about the principal's IC constraint

(a) Assume that at the boundary only the agent's truth-telling constraint is binding, and we are in the lowest possible state. I will show that introducing a positive $\Delta U_A(L, \omega)$ to this equilibrium strictly improves the principal's payoff. To this end, recall that the promise-keeping and truth-telling constraints are given by

$$\begin{aligned} U_A(\omega)(1 - \delta_A) &= u_A(\omega) + \delta_A(1 - p)\Delta U_A(L, \omega) \\ \delta_A \Delta U_A(L, \omega) + \delta_A \Pr(A|H, \omega)U_A(\omega) &= (\Pr(A|H, \omega) - \Pr(A|L, \omega))\theta_L. \end{aligned}$$

Implicitly differentiating to promise-keeping constraint with respect to $\Delta U_A(L, \omega)$ gives us

$$p\theta_H \frac{\partial x_H}{\partial \Delta U_A(L, \omega)} + (1 - p)\theta_L \frac{\partial x_L}{\partial \Delta U_A(L, \omega)} + \delta_A(1 - p) = 0,$$

and repeating the same for the truth-telling constraint gives us

$$\frac{\partial x_L}{\partial \Delta U_A(L, \omega)}\theta_L = \frac{\partial x_H}{\partial \Delta U_A(L, \omega)}\theta_L - \delta_A.$$

Solving the equations together gives us

$$\frac{\partial x_H}{\partial \Delta U_A(L, \omega)} = 0 \quad \text{and} \quad \frac{\partial x_L}{\partial \Delta U_A(L, \omega)} = -\frac{\delta_A}{\theta_L}.$$

Finally, note that because we are in a scenario where the agent's truth-telling constraint is binding, $x_L > 0$. Thus, introducing a positive likelihood of moving away from the state allows us to lower the level of favoritism. The principal's payoff is then unambiguously improved since $\frac{\partial \Delta U_P(L, \omega)}{\partial \Delta U_A(L, \omega)} > 0$ and $\frac{\partial \Delta U_P(L, \omega)}{\partial x_L} < 0$. Note that this also implies that an absorbing state cannot exist if both the agent's and the principal's constraints are binding. Finally, for low states there is no situation where only the principal's IC could be binding in the absorbing state.

(b) Now, assume that the agent's truth-telling constraint is binding at the highest possible state, and only consider adding a penalty for rejecting a high-quality project. The two constraints are given by

$$\begin{aligned} U_A(\omega)(1 - \delta_A) &= u_A(\omega) + \delta_A p \Pr(R|H, \omega) \Delta U_A(R|H, \omega) \\ -\delta_A \Pr(R|H, \omega) \Delta U_A(R|H, \omega) + \delta_A \Pr(A|H, \omega)U_A(\omega) &= (\Pr(A|H, \omega) - \Pr(A|L, \omega))\theta_L, \end{aligned}$$

and implicitly differentiating the promise-keeping constraint gives us (for $\Delta U_A(R|H, \omega) \rightarrow 0$)

$$p\theta_H \frac{\partial x_H}{\partial \Delta U_A(R|H, \omega)} + (1 - p)\theta_L \frac{\partial x_L}{\partial \Delta U_A(R|H, \omega)} + \delta_A p \Pr(R|H, \omega) = 0,$$

and for the truth-telling constraint, we have (for $\Delta U_A(R|H, \omega) \rightarrow 0$)

$$\frac{\partial x_L}{\partial \Delta U_A(R|H, \omega)} \theta_L = \left(\delta_A \Pr(R|H, \omega) + \frac{\partial x_H}{\partial \Delta U_A(R|H, \omega)} (\theta_L - \delta_A U_A(\omega)) \right).$$

Solving the conditions together gives us then

$$\begin{aligned} \frac{\partial x_H}{\partial \Delta U_A(R|H, \omega)} &= -\frac{\delta_A \Pr(R|H, \omega)}{(E(\theta) - (1-p)\delta_A U_A(\omega))} < 0 \\ \frac{\partial x_L}{\partial \Delta U_A(R|H, \omega)} \theta_L &= \delta_A \Pr(R|H, \omega) \left(\frac{p(\theta_H - \theta_L) + p\delta_A U_A(\omega)}{(E(\theta) - (1-p)\delta_A U_A(\omega))} \right) > 0, \end{aligned}$$

where $(E(\theta) - (1-p)\delta_A U_A(\omega)) > 0$ follows from the fact that, near the stationary equilibrium, $\Delta U_A(R|H, \omega) = 0$, which implies that

$$0 = \frac{U_A(\omega)[1 - \delta_A(1 - (1-p)\Pr(A|H, \omega))] - E(\theta)(\Pr(A|H, \omega) - E(\theta))}{\delta_A \Pr(R|H, \omega)},$$

which gives us $U_A(\omega) = \frac{E(\theta)(\Pr(A|H, \omega) - E(\theta))}{[1 - \delta_A(1 - (1-p)\Pr(A|H, \omega))]}$ and so

$$\begin{aligned} (E(\theta) - (1-p)\delta_A U_A(\omega)) &= \\ E(\theta) \left(1 - (1-p)\delta_A \frac{(\Pr(A|H, \omega) - E(\theta))}{[1 - \delta_A(1 - (1-p)\Pr(A|H, \omega))]} \right) &= E(\theta) \left(\frac{1 - \delta_A + (1-p)\delta_A E(\theta)}{[1 - \delta_A(1 - (1-p)\Pr(A|H, \omega))]} \right) > 0. \end{aligned}$$

Note that a penalty means decreasing $\Delta U_A(R|H, \omega)$, so that the signs of the derivatives are reversed. And so decreasing $\Delta U_A(R|H, \omega)$ allows us to decrease x_L and increase x_H , which will improve the principal's payoff, both through the more efficient stage-game payoff (agent-IC doesn't allow sufficient spread in the distortions) and return to a more favorable equilibrium. So stationarity cannot be optimal.

(c) Finally, suppose that we are in the region where only the principal's IC constraint is binding. Now, from the agent's side we only have the promise-keeping constraint,

$$\begin{aligned} U_A(\omega)(1 - \delta_A) &= u_A(\omega) + \delta_A(1 - p)\Delta U_A(L, \omega) \\ &\quad + \delta_A p [\Pr(R|H, \omega)\Delta U_A(R|H, \omega) + \Pr(A|H, \omega)\Delta U_A(A|H, \omega)] \end{aligned}$$

and from the principal's side we have

$$x_i \leq \delta(U_P(\omega) + \Delta U_P(A|i, \omega))$$

and

$$\begin{aligned} U_P(\omega)(1 - \delta_P) &= u_P(\omega) + \delta_P(1 - p)\Delta U_P(L, \omega) \\ &\quad + \delta_P p [\Pr(R|H, \omega)\Delta U_P(R|H, \omega) + \Pr(A|H, \omega)\Delta U_P(A|H, \omega)]. \end{aligned}$$

Now, to establish the result we need to first bound the derivative for $\frac{\partial \Delta U_P(A|H, \omega)}{\Delta U_A(A|H, \omega)}$. But this slope is bounded by the stationary equilibrium, since we want to only show an improvement to a stationary equilibrium. Thus, we have for the agent that

$$(1 - \delta_A) U_A(\omega) = u_A(\omega).$$

So, to increase his payoff, we have that

$$(1 - p) \theta_L \frac{\partial x_L}{\partial u_A(\omega)} = 1 - p \theta_H \frac{\partial x_H}{\partial u_A(\omega)}.$$

For the principal, we know from the characteristics of the optimal stationary solution that it is the favoritism in the high state that is binding (efficient distortion in a favorable state implies $x_H > x_L$). For this constraint, we have

$$\frac{\partial x_H}{\partial u_A(\omega)} - \frac{\delta_P}{1 - \delta_P} \left(-p x_H \frac{\partial x_H}{\partial u_A(\omega)} - (1 - p) x_L \frac{\partial x_L}{\partial u_A(\omega)} \right) = 0,$$

which we can combine with the agent's promise-keeping constraint to yield

$$\frac{\partial x_H}{\partial u_A(\omega)} = - \frac{\delta_P x_L \frac{1}{\theta_L}}{\left[(1 - \delta_P) + p \delta_P x_L \left(\frac{x_H}{x_L} - \frac{\theta_H}{\theta_L} \right) \right]}.$$

Then, we can write the effect on the principal's payoff as

$$\begin{aligned} \frac{\partial u_P(\omega)}{\partial u_A(\omega)} &= + \left(x_L \frac{1}{\theta_L} p \theta_H - p x_H \right) \frac{\partial x_H}{\partial u_A(\omega)} - \frac{x_L}{\theta_L} \\ &= - \left(\frac{\frac{x_L}{\theta_L} (1 - \delta_P)}{\left[(1 - \delta_P) - p \delta_P x_L \left(\frac{\theta_H}{\theta_L} - \frac{x_H}{x_L} \right) \right]} \right), \end{aligned}$$

which gives us the slope of the stationary frontier, and we can consider a small move along the frontier from the current state. Now, suppose we lower the agent's continuation value following acceptance. We have $\Delta U_A(A|H, \omega) < 0$ and consider a small change. The current value is given by

$$U_A(\omega) (1 - \delta_A) = u_A(\omega) + \delta_A p \Pr(A|H, \omega) \Delta U_A(A|H, \omega),$$

from which we get that for small changes it needs to be that

$$p \theta_H \frac{\partial x_H}{\partial \Delta U_A(A|H, \omega)} + (1 - p) \theta_L \frac{\partial x_L}{\partial \Delta U_A(A|H, \omega)} + \delta_A p \Pr(A|H, \omega) = 0,$$

and for the principal, the IC constraint must remain binding, so that

$$x_H - \delta_P (U_P(\omega) + \Delta U_P(A|H, \omega)) = 0,$$

while the payoff identity gives us

$$U_P(\omega) = \frac{u_P(\omega) + \delta_P p [\Pr(A|H, \omega) \Delta U_P(A|H, \omega)]}{(1 - \delta_P)},$$

so that we have

$$U_P(\omega) + \Delta U_P(A|H, \omega) = \frac{u_P(\omega) + [(1 - \delta_P) + \delta_P p \Pr(A|H, \omega)] \Delta U_P(A|H, \omega)}{(1 - \delta_P)}.$$

Then, implicitly differentiating the IC constraint gives us

$$\frac{\partial x_H}{\partial \Delta U_A(A|H, \omega)} - \frac{\delta_P}{(1-\delta_P)} \left(\frac{-px_H \frac{\partial x_H}{\partial \Delta U_A(A|H, \omega)} - (1-p)x_L \frac{\partial x_L}{\partial \Delta U_A(A|H, \omega)}}{+ [(1-\delta_P) + \delta_P p \Pr(A|H, \omega)] \frac{\Delta U_P(A|H, \omega)}{\Delta U_A(A|H, \omega)}} \right) = 0,$$

and from the agent's promise-keeping constraint we get, again implicitly differentiating, that

$$(1-p) \frac{\partial x_L}{\partial \Delta U_A(A|H, \omega)} = - \left(\frac{\delta_{AP}}{\theta_L} \Pr(A|H, \omega) + p \frac{\theta_H}{\theta_L} \frac{\partial x_H}{\partial \Delta U_A(A|H, \omega)} \right),$$

so that the principal's IC constraint finally simplifies to

$$\frac{dx_H}{d\Delta U_A(A|H, \omega)} = \frac{\left(\delta_P x_L \frac{\delta_{AP}}{\theta_L} \Pr(A|H, \omega) + \delta_P [(1-\delta_P) + \delta_P p \Pr(A|H, \omega)] \frac{\Delta U_P(A|H, \omega)}{\Delta U_A(A|H, \omega)} \right)}{\left[(1-\delta_P) - p \delta_P x_L \left(\frac{\theta_H}{\theta_L} - \frac{x_H}{x_L} \right) \right]}.$$

Now, we need $\frac{dx_H}{d\Delta U_A(A|H, \omega)} < 0$ for us to benefit from introducing a penalty for acceptance, which requires that

$$\delta_P x_L \frac{\delta_{AP}}{\theta_L} \Pr(A|H, \omega) + \delta_P [(1-\delta_P) + \delta_P p \Pr(A|H, \omega)] \frac{\Delta U_P(A|H, \omega)}{\Delta U_A(A|H, \omega)} < 0,$$

which becomes, after substituting $\frac{\Delta U_P(A|H, \omega)}{\Delta U_A(A|H, \omega)}$ from the stationary frontier

$$- \Pr(A|H, \omega) \delta_P p \left[\frac{p \delta_P x_L \left(\frac{\theta_H}{\theta_L} - \frac{x_H}{x_L} \right)}{\left[(1-\delta_P) - p \delta_P x_L \left(\frac{\theta_H}{\theta_L} - \frac{x_H}{x_L} \right) \right]} \right] < \left(\frac{(1-\delta_P)^2}{\left[(1-\delta_P) - p \delta_P x_L \left(\frac{\theta_H}{\theta_L} - \frac{x_H}{x_L} \right) \right]} \right).$$

And since the right-hand side is positive while the left-hand side is negative, we have established the result. The last check would be to take the end-point where both high and low states are binding. But that follows similarly since we are increasing the spread and it is the high state that remains binding while the low state is relaxed. So we have established that the equilibrium will not have an absorbing state, with the potential exception of actually reaching the frontier.

Lemma 9 *Around the state that maximizes the principal's payoff, the favoritism towards the low-quality project is positive ($x_{L, \omega_0} > 0$), the treatment of high-quality projects is ambiguous ($x_{L, \omega_0} \gtrless 0$), the admittance of mediocre projects is rewarded ($\Delta U_A(L, \omega_0) > 0$), while the acceptance of high-quality projects is rewarded ($\Delta U_A(A|H, \omega_0) > 0$) and the rejection of high-quality projects is punished ($\Delta U_A(R|H, \omega_0) < 0$).*

The principal's payoff in a given state is given by

$$U_P(\omega) (1 - \delta_P) = u_P(\omega) + \delta_P (1 - p) \Delta U_P(L, \omega) + \delta_P p [\Pr(R|H, \omega) \Delta U_P(R|H, \omega) + \Pr(A|H, \omega) \Delta U_P(A|H, \omega)].$$

The principle of optimality allows us to consider the optimal choice of variables in any given state, conditional on optimal play in all the other states, observation that we will use repeatedly below without repeating the logic. For $x_{L, \omega}$, consider the optimal choice of $\Delta U_A(L, \omega)$, with the relevant constraints given by equations 33 and 34 and the associated derivatives. The only variable that is directly linked to $\Delta U_A(L, \omega)$ is $x_{L, \omega}$ and the optimal choice thus solves

$$-(1-p)x_{L,\omega}\frac{\partial x_{L,\omega}}{\partial \Delta U_A(L,\omega)} + \delta_P(1-p)\frac{\partial \Delta U_P(L,\omega)}{\partial \Delta U_A(L,\omega)} = 0,$$

and substituting in $\frac{\partial x_{L,\omega}}{\partial \Delta U_A(L,\omega)}$ and using $\delta_P = \delta_A$ gives us

$$x_{L,\omega} = -\theta_L \frac{\partial \Delta U_P(L,\omega)}{\partial \Delta U_A(L,\omega)},$$

which will be positive around the maximum, where $\frac{\partial \Delta U_P(L,\omega)}{\partial \Delta U_A(L,\omega)} < 0$.

For $x_{H,\omega}$, we can repeat the exercise and the first-order condition and obtain

$$\begin{aligned} U_P(\omega)(1-\delta_P) &= -px_{H,\omega} + \delta_P(1-p)\frac{\partial \Delta U_P(L,\omega)}{\partial \Delta U_A(L,\omega)}\frac{\partial \Delta U_A(L,\omega)}{\partial x_{H,\omega}} \\ &\quad + \delta_PP[-\Delta U_P(R|H,\omega)] + \delta_PP\left[\Pr(R|H,\omega)\frac{\partial \Delta U_P(R|H,\omega)}{\partial \Delta U_A(R|H,\omega)}\frac{\partial \Delta U_A(R|H,\omega)}{\partial x_{H,\omega}}\right] + \delta_PP[\Delta U_P(A|H,\omega)] = \\ &= 0. \end{aligned}$$

Using the condition for the optimal choice of $x_{L,\omega}$ and $\frac{\partial \Delta U_A(L,\omega)}{\partial x_{H,\omega}}$ and $\frac{\partial \Delta U_A(R|H,\omega)}{\partial x_{H,\omega}}$, we can write the above as

$$\begin{aligned} x_{H,\omega} &= (1-p)\frac{x_{L,\omega}(\delta_A(U_A(\omega) + \Delta U_A(A|H,\omega)) + (\theta_H - \theta_L))}{\theta_L} + \delta_P(\Delta U_P(A|H,\omega) - \Delta U_P(R|H,\omega)) \\ &\quad - \left(\frac{E(\theta) - E(\theta)^2 - U_A(\omega) + \delta_A p(U_A(\omega) + \Delta U_A(A|H,\omega))}{\Pr(R|H,\omega)}\right)\frac{\partial \Delta U_P(R|H,\omega)}{\partial \Delta U_A(R|H,\omega)}. \end{aligned}$$

This result is ambiguous because using $x_{H,\omega} > 0$ allows us to balance the compensation delivered through $x_{L,\omega}$, but it also changes the probability of the two continuation outcomes, $(\Delta U_P(A|H,\omega), \Delta U_P(R|H,\omega))$ where acceptance may yield a lower payoff reduction on the principal than rejection, and most importantly, increasing $x_{H,\omega}$ requires increasing the severity of penalty through $\Delta U_A(R|H,\omega)$ and this is increasingly costly to the principal, which may induce a negative $x_{H,\omega}$.

The optimal choice of $\Delta U_A(A|H,\omega)$ solves

$$\begin{aligned} (1-p)\frac{\partial \Delta U_P(L,\omega)}{\partial \Delta U_A(L,\omega)}\frac{\partial \Delta U_A(L,\omega)}{\partial \Delta U_A(A|H,\omega)} \\ + p\left[\Pr(R|H,\omega)\frac{\partial \Delta U_P(R|H,\omega)}{\partial \Delta U_A(R|H,\omega)}\frac{\partial \Delta U_A(R|H,\omega)}{\partial \Delta U_A(A|H,\omega)} + \Pr(A|H,\omega)\frac{\partial \Delta U_P(A|H,\omega)}{\partial \Delta U_A(A|H,\omega)}\right] = 0, \end{aligned}$$

and again using the derivatives from above, this simplifies to

$$\frac{\partial \Delta U_P(A|H,\omega)}{\partial \Delta U_A(A|H,\omega)} = (1-p)\frac{\partial \Delta U_P(L,\omega)}{\partial \Delta U_A(L,\omega)} + p\frac{\partial \Delta U_P(R|H,\omega)}{\partial \Delta U_A(R|H,\omega)}.$$

Finally, while $U_A(\omega)$ is determined through the transitions for the remainder of the game, the first-period choice is a choice for the principal and given by

$$\delta_P(1-p)\frac{\partial \Delta U_P(L,\omega)}{\partial \Delta U_A(L,\omega)}\frac{\partial \Delta U_A(L,\omega)}{\partial U_A(\omega)} + \delta_PP\left[\Pr(R|H,\omega)\frac{\partial \Delta U_P(R|H,\omega)}{\partial \Delta U_A(R|H,\omega)}\frac{\partial \Delta U_A(R|H,\omega)}{\partial U_A(\omega)}\right] = 0,$$

which then simplifies to

$$\frac{\partial \Delta U_P(R|H,\omega)}{\partial \Delta U_A(R|H,\omega)} = -\frac{\partial \Delta U_P(L,\omega)}{\partial \Delta U_A(L,\omega)}\frac{(1-p)[1-\delta_A(1+p\Pr(A|H,\omega))]}{p[1-\delta_A(1-(1-p)\Pr(A|H,\omega))]},$$

which implies that at the initial optimum, we do have a penalty for rejection and a reward for admitting a mediocre project (the slopes need to be of opposite signs). Further, we can use this to pin down the initial effect of accepting high-quality projects, after substituting the relationship into the first-order condition for $\Delta U_A(A|H, \omega)$ as

$$\frac{\partial \Delta U_P(A|H, \omega)}{\partial \Delta U_A(A|H, \omega)} = (1-p) \frac{\partial \Delta U_P(L, \omega)}{\partial \Delta U_A(L, \omega)} \left(\frac{\delta_A \Pr(A|H, \omega)}{[1 - \delta_A(1 - (1-p) \Pr(A|H, \omega))]} \right) < 0,$$

which implies an initial reward for the acceptance of high-quality projects.

Lemma 10 $x_{L, \omega}$ will be monotone increasing in the state

First, observe that $\frac{\partial \Delta U_A(L, \omega)}{\partial U_A(\omega)} = \frac{[1 - \delta_A(1 + p \Pr(A|H, \omega))]}{\delta_A} \geq -1$. This implies that $\frac{\partial U_A(L, \omega)}{\partial U_A(\omega)} > 0$, holding all other variables constant. This result implies that an increase in $U_A(\omega)$ decreases $\frac{\partial \Delta U_P(L, \omega)}{\partial \Delta U_A(L, \omega)}$, holding other things constant. This, in turn, implies that, for the original values,

$$x_{L, \omega} < -\theta_L \frac{\partial \Delta U_P(L, \omega)}{\partial \Delta U_A(L, \omega)},$$

and since $\frac{\partial \Delta U_A(L, \omega)}{\partial x_L} = -\frac{\theta_L}{\delta_A} < 0$, the new optimal solution involves an increase in $x_{L, \omega}$, which in turn involves an increase in $\frac{\partial \Delta U_P(L, \omega)}{\partial \Delta U_A(L, \omega)}$ until equality is restored.

Conjecture 11 $x_{H, \omega}$ will be monotone increasing in the state

From above, we have that the agent's first-order condition is given by

$$x_{H, \omega} = (1-p) \left[\frac{x_{L, \omega}(\delta_A U_A(A|H, \omega) + (\theta_H - \theta_L))}{\theta_L} \right] + \delta_P (\Delta U_P(A|H, \omega) - \Delta U_P(R|H, \omega)) \\ - \left(\frac{E(\theta) - E(\theta)^2 - U_A(\omega) + \delta_A p U_A(A|H, \omega)}{\Pr(R|H, \omega)} \right) \frac{\partial \Delta U_P(R|H, \omega)}{\partial \Delta U_A(R|H, \omega)}.$$

Holding $x_{H, \omega}$ and $U_A(A|H, \omega)$ constant, the effect on the RHS then

$$(1-p) \left[\frac{x_{L, \omega}(\delta_A U_A(A|H, \omega) + (\theta_H - \theta_L))}{\theta_L} \right] \frac{\partial x_{L, \omega}}{\partial U_A(\omega)} \\ + \frac{1}{\Pr(R|H, \omega)} [\delta_A (1 - (1-p) \Pr(A|H, \omega))] \frac{\partial \Delta U_P(R|H, \omega)}{\partial \Delta U_A(R|H, \omega)} \\ - \left(\frac{E(\theta) - E(\theta)^2 - U_A(\omega) + \delta_A p U_A(A|H, \omega)}{\Pr(R|H, \omega)} \right) \frac{\partial \left(\frac{\partial \Delta U_P(R|H, \omega)}{\partial \Delta U_A(R|H, \omega)} \right)}{\partial U_A(\omega)}.$$

The first line is positive due to $\frac{\partial x_{L, \omega}}{\partial U_A(\omega)} > 0$, driven by the fact that increasing compensation through $x_{H, \omega}$ balances the need to increase compensation through $x_{L, \omega}$. The second line is the net effect of the change in the location of the optimal penalty, and as long as a penalty brings us to the upward-sloping portion of the payoff frontier (which we would expect but have not yet established since increasing the penalty allows us to lower the reward while improving the principal's payoff). And

finally, because $\frac{\partial \Delta U_A(R|H, \omega)}{\partial U_A(\omega)} > 0$, the concavity of the frontier implies that $\frac{\partial \left(\frac{\partial \Delta U_P(R|H, \omega)}{\partial \Delta U_A(R|H, \omega)} \right)}{\partial U_A(\omega)} < 0$ and thus the third effect is positive as well. As a result, bringing the expression back to balance requires us to increase $x_{H, \omega}$ (with the challenge that third component is becoming increasingly positive as we change $x_{H, \omega}$ as the probability of a penalty is correspondingly decreasing when we increase the level of favoritism).

Conjecture 12 *The variance in the continuation values needed to sustain truthful communication is decreasing in ω .*

Holding the other components constant, we have that the incentives provided through the continuation game to the agent are given by

$$\Delta U_A(L, \omega) - \Pr(R|H, \omega) \Delta U_A(R|H, \omega),$$

and taking the derivative with respect to the current state, we get

$$\begin{aligned} & \frac{\partial [\Delta U_A(L, \omega) - \Pr(R|H, \omega) \Delta U_A(R|H, \omega)]}{\partial U_A(\omega)} \\ &= \left(\frac{[1 - \delta_A(1 + p \Pr(A|H, \omega))]}{\delta_A} - \frac{[1 - \delta_A(1 - (1 - p) \Pr(A|H, \omega))]}{\delta_A} \right) = -\Pr(A|H, \omega) < 0. \end{aligned}$$

Now, there are additional adjustments that occur to this through changes in $x_{L, \omega}$ and $x_{H, \omega}$. We know that increasing $x_{L, \omega}$ allows us to decrease $\Delta U_A(L, \omega)$ even further but an increase in $x_{H, \omega}$ does push $\Delta U_A(R|H, \omega)$ back out, and we need to make sure that this change doesn't exceed the original adjustment. The economic intuition is that in highly favorable states the agent directly values the relationship more and thus smaller variation in continuation values is needed to sustain informative communication.

Conjecture 13 *Mediocre projects receive positive favoritism in all states, with $x_{L, \omega} > 0$. Further, the return from $\omega < \omega_0$ through the revelation of a mediocre project is always rewarded by a return to a state above ω_0 .*

The first-order condition for the choice of $x_{L, \omega}$ is $x_L = -\theta_L \frac{\partial \Delta U_P(L, \omega)}{\partial \Delta U_A(L, \omega)}$. If the solution was negative, this would imply that $\frac{\partial \Delta U_P(L, \omega)}{\partial \Delta U_A(L, \omega)} > 0$. But then it would make sense to readjust parameters in a way that the reward can be higher while the current penalty is stronger through other means. so the only reason for the game not to return to a state that is favorable to the principal is if the penalty is sufficiently large so that the return to a more favorable state for the agent needs to be optimally delayed.

B Continuous-state model

This Appendix considers the additional insights that can be obtained when both the principal's implementation cost and the agent's project quality are continuously distributed on $[0,1]$, with cumulative distribution functions $F_P(c)$ and $F_A(\theta_i)$ (which, for some explicit illustrations, I will assume to be uniform). Given the complexity of the potential dynamics, no tractable solution exists for the general problem. Instead, I will characterize the optimal stationary equilibrium, to illustrate how the richer state space helps to illustrate the additional distortions present that are not present in the two-state example. The principal's payoff continues to be $u_P = \theta_i - c$, while the agent's payoff is given by $u_A(\theta_i) = b(1 - \alpha) + \alpha\theta_i$. Finally, I will focus only on the truth-telling constraint of the agent. Otherwise, the game behaves like discussed earlier. The agent learns θ_i , sends a message m_i to the principal, after which c is publicly observed and the principal makes the implementation decision according to $\Pr(A|m_i, c)$. The goal is to characterize the optimal $\Pr(A|m_i, c)$. Finally, for simplicity, I assume that the principal's cost distribution is such that the inverse hazard rate $\frac{1-F_P(c)}{f_P(c)}$ is monotone decreasing.

B.1 Attaining the first-best

Let $\Pr(A|m_i) = \int \Pr(A|m_i, c) dF_P(c) dc$ denote the expected probability of acceptance following a given message. ^c Then, assuming that we can attain the first-best, the continuation equilibrium is stationary, and telling the truth gives an expected payoff of

$$\Pr(A|\theta_i) u_A(\theta_i) + \delta_A V_A^{cont}. \quad (35)$$

Similarly, if the agent chooses to lie and sends a message $m_i = \tilde{\theta}_i$, his payoff is given by

$$\Pr(A|\tilde{\theta}_i) u_A(\theta_i) + \delta_A \left[\left(1 - \Pr(A|\tilde{\theta}_i)\right) V_A^{cont} + \Pr(A|\tilde{\theta}_i) V_A^{dev} \right]. \quad (36)$$

The optimal decision rule involves the principal accepting the project whenever $\theta_i \geq c$, so that $\Pr(A|\tilde{\theta}_i) = F_P(\tilde{\theta}_i)$. Then, the truth-telling constraint reduces to, for all θ_i and $\tilde{\theta}_i$, to

$$V_A^{cont} - V_A^{dev} \geq \frac{1}{\delta_A} \left(1 - \frac{F_P(\theta_i)}{F_P(\tilde{\theta}_i)} \right) u_A(\theta_i). \quad (37)$$

Two observations follow. First, the optimal deviation for the agent is to always send the maximal message, which, given the assumed shared support and the first-best decision rule, leads to acceptance with probability 1. Intuitively, if it is optimal for the agent to risk burning his reputation, he should go all in. Second, as long as the agent cares enough about the quality of the project, the

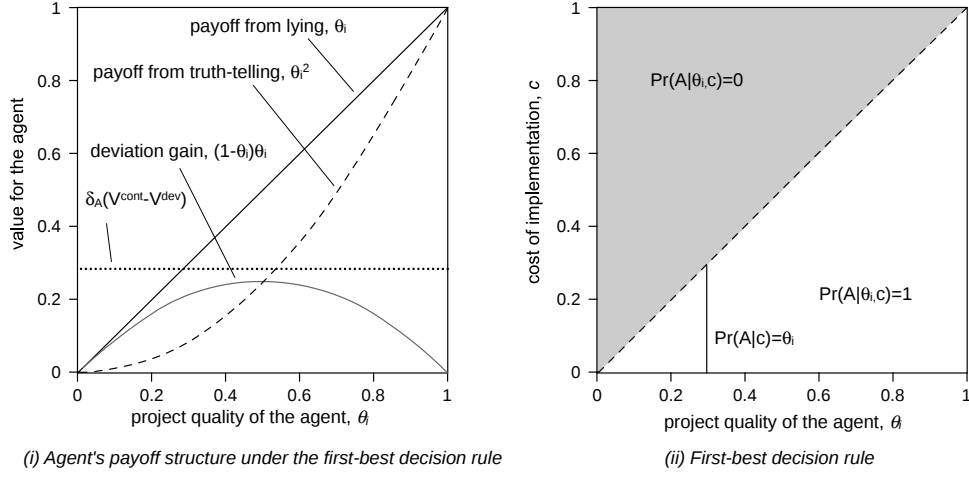


Figure 11: Sustainable truth-telling under the first-best decision rule.

deviation temptation is maximized for interior θ_i . Since $F_P(\tilde{\theta}_i) = 1$, the deviation temptation is maximized to θ_i for which

$$\theta_i = \frac{(1 - F_P(\theta_i))}{f_P(\theta_i)} - \frac{b(1 - \alpha)}{\alpha}. \quad (38)$$

In the case of the uniform distribution, this yields $\theta_i = \frac{1}{2} - \frac{b(1-\alpha)}{2\alpha} \leq \frac{1}{2}$. Intuitively, when α is sufficiently large and the state is low enough, the return to pushing for acceptance is not worth the destruction of the reputation when caught lying. Similarly, when the state is high enough, then the project is relatively likely to be accepted even without exaggeration, and the small improvement in the probability of acceptance is not worth risking the reputation. It is this non-monotonicity of the deviation temptation that is the main difference to the two-state variant. The more the agent cares only about acceptance, the more the maximal deviation temptation gets skewed towards low-quality projects, as the weight on increasing the likelihood of acceptance relative to the value the agent receives conditional on acceptance, increases.

This first-best solution is illustrated in Figure 11 for the uniform distribution and $b = 1$, so that acceptance alone creates no value to the agent. Panel (ii) plots the first-best decision rule. Simply, the agent tells the truth and the principal implements the agent's proposal if the cost is below the value of the project. Conditional on the decision rule, panel (i) illustrates the state-contingent payoffs of the agent, highlighting how the deviation temptation is maximized at $\theta_i = \frac{1}{2}$. Thus, if the truth-telling constraint is satisfied for $\theta_i = \frac{1}{2}$, it is slack for all other states. A quick valuation of the constraint reveals that truth-telling and first-best decision rule are incentive-compatible as long as $\delta_A \geq \frac{3}{4}$.

For general preferences, the condition becomes

$$\frac{\delta_A}{1 - \delta_A} \int_{\theta_i} (F_P(\theta_i) - F_P(E(\theta_i))) \theta_i dF_A(\theta_i) \geq \left[\frac{(1 - F_P(\hat{\theta}_i))^2}{f_P(\hat{\theta}_i)} \right], \quad (39)$$

where the left-hand side captures the value of the truthful relationship, while the right-hand side captures the gain from deviating, with $\hat{\theta}_i$ solving equation 38. Note that while the left-hand side of the expression is independent of the agent's preferences (α, b) , they influence the deviation temptation on the right-hand side. In particular, decreasing the agent's payoff-sensitivity to the state (α) and increasing his reward for acceptance (b) both increase his reneging temptation by pushing down $\hat{\theta}_i$. For the uniform distribution, we can explicitly solve for the threshold, which is given by

$$\delta_A \geq \frac{3(\alpha + b(1 - \alpha))^2}{(\alpha^2 + 3(\alpha + b(1 - \alpha))^2)}, \quad (40)$$

with $\frac{\partial \delta_A}{\partial \alpha} < 0$ and $\frac{\partial \delta_A}{\partial b} > 0$. Intuitively, the more the agent cares (in relative terms) about the quality of the implemented project, the more valuable the relationship and, as a result, the less patient the agent needs to be and still achieve the first-best solution.

B.2 Managing the truth-telling constraint

If the agent is not sufficiently patient, then the first-best decision rule is unable to sustain truth-telling by the agent. In this section, I will consider how to optimally alter the decision rule to sustain truth-telling by the agent (because we assume that the principal can commit to the decision rule, the revelation principle applies and we can focus on truthful mechanisms). From above, recall that the truth-telling constraint was given by

$$V_A^{cont} - V_A^{dev} \geq \frac{1}{\delta_A} \left(1 - \frac{\Pr(A|\theta_i)}{\Pr(A|\tilde{\theta}_i)} \right) u_A(\theta_i). \quad (41)$$

The expression reveals immediately the three avenues through which the principal can restore incentive-compatibility: V_A^{cont} , $\Pr(A|\tilde{\theta}_i)$ and $\Pr(A|\theta_i)$. I will discuss the optimal use of these three avenues separately to reveal the economic logic behind them:

1. General favoritism: The first immediate means of managing the agent's incentives to deviate is to increase the agent's influence in the stage game. In other words, we can increase $E[\Pr(A|\theta_i)]$ to increase the agent's expected payoff and thus the value of the relationship, V_A^{cont} . The question is then what is the most efficient means of increasing the agent's stage-game payoff, ignoring any other constraints, and the solution is given by the following lemma:

Lemma 14 *The least-cost means of increasing the agent's continuation value is to use a linear distortion in the acceptance rule, with the principal implementing the project whenever $c \leq \bar{c}(\theta_i)$, where*

$$\bar{c}(\theta_i) = \bar{c}(0) + \theta_i \left(1 + \frac{[\alpha \bar{c}(0)]}{b(1-\alpha)} \right) \geq 0. \quad (42)$$

As $\alpha \rightarrow 1$, $\bar{c}(0) \rightarrow 0$ and $\bar{c}(\theta_i) = (1 + \eta) \theta_i$, for $\eta \geq 0$.

Proof. See Appendix B.3.1 ■

In other words, because the agent's payoff is linear to the value of the idea, the least distortionary means of delivering a given utility to the agent is a linear distortion, where the principal is willing to implement projects at a loss, but those losses are spread across the states this way. In particular, as long as $\alpha > 0$, the principal can deliver more value by favoring the high-quality proposals relatively more. As $\alpha \rightarrow 1$, the distortion becomes proportional to the value of the project, while if $\alpha \rightarrow 0$, the distortion becomes additive, where the principal introduces the same absolute level of favoritism for all states. Note that because the probability distributions are common knowledge, this result holds independent of the underlying distributions.

2. Discrimination against the best proposals: The second observation that follows from equation 41 is that the agent's optimal deviation is going to be to the proposal that maximizes the probability of acceptance., and the deviation temptation is increasing in the maximum probability of acceptance. Thus, the second means of satisfying the agent's truth-telling constraint is to lower the maximum probability of acceptance, which means that the agent's best proposals will be discriminated against. Noting that the cost of distortions is increasing as we move away from the diagonal, the least-cost way of lowering the deviation temptation is by introducing a hard cap $\bar{\theta}_i$, where proposals above this quality are implemented with a common probability $\Pr(A|\bar{\theta}_i) = \bar{p} < 1$. Further, combining (1) and (2), the transition from general favoritism into discrimination must be continuous, so the kink is given by the point at which $F_P \left(\bar{c}(0) + \bar{\theta}_i \left(1 + \frac{[\alpha \bar{c}(0)]}{b(1-\alpha)} \right) \right) = \bar{p}(\alpha, b)$.

To understand why the maximal acceptance probability plays a role even if the deviation temptation is maximized for intermediate types is as follows. When the agent chooses whether to lie or not, he doesn't yet know whether the lie is actually needed to induce acceptance. The lie is needed only when the cost of implementation is high enough, while truth would be enough when the cost is low enough. When choosing whether to mislead the principal, the agent is balancing these two forces. Now, if we lower the maximal probability of acceptance, so that the principal never implements the agent's proposal when her cost is high enough, the relative efficiency of the lie is lowered: it is more likely that telling the truth would have been sufficient to induce acceptance. Thus, exaggeration becomes less attractive.

In relation to the two-state framework of the main analysis, the level of distortion for the better project in the two-state framework blends these two effects (increasing the continuation value and decreasing the reneging temptation), where the direction of forces works in the opposite directions.

In the continuous-state model, they become more decoupled as the principal can discriminate against the best projects while still providing favoritism towards above-average quality projects.

3. Leniency towards average proposals: As noted above, general favoritism is expensive, in the sense that it increases the agent's payoff for all types, even for those that would be willing to be truthful under a less favorable decision rule. An alternative means is then to increase the acceptance probability only for those types for whom it is directly needed. The basic tradeoff is that such focused leniency will be less likely to be needed (because it is needed only for some states) but it will be more expensive when needed (since the distortion will need to be higher).

Following the logic from above, the deviation temptation is maximized for intermediate states. Then, the concavity of the deviation temptation implies that if the truth-telling constraint of equation 41 is violated at some $\hat{\theta}_i$, then there exist bounds $\underline{\theta} < \hat{\theta}_i < \bar{\theta}$ for which the deviation temptation is exactly satisfied. Then, letting $\bar{c}(\theta_i|\alpha, b)$ denote the (linear) general degree of favoritism, and $\bar{p}(\alpha, b)$ the degree of discrimination against the best proposals, we can write the truth-telling constraint as

$$V_A^{cont} - V_A^{dev} \geq \frac{1}{\delta_A} \left(1 - \frac{\Pr(A|\theta_i)}{\bar{p}(\alpha, b)} \right) u_A(\theta_i), \quad (43)$$

which then, for $\underline{\theta}$, satisfies exactly

$$V_A^{cont} - V_A^{dev} = \frac{1}{\delta_A} \left(1 - \frac{F_P(\bar{c}(\underline{\theta}_i|\alpha, b))}{\bar{p}(\alpha, b)} \right) u_A(\underline{\theta}_i), \quad (44)$$

which allows us to define the level of distortion for the intermediate ranges as

$$\Pr(A|\theta_i) = \bar{p}(\alpha, b) - (\bar{p}(\alpha, b) - F_P(\bar{c}(\underline{\theta}_i|\alpha, b))) \frac{u_A(\underline{\theta}_i)}{u_A(\theta_i)} > F_P(\bar{c}(\theta_i|\alpha, b)). \quad (45)$$

As shown in Appendix B.3.2, we can then write, for the uniform distribution and $\alpha = 1$, the boundaries as

$$\{\underline{\theta}, \bar{\theta}\} = \frac{\bar{p}}{2(1+\eta)} (1 \pm X), \quad X = \sqrt{1 - \frac{4(1+\eta)\delta_A\Delta V^{cont}}{\bar{p}}}, \quad (46)$$

where ΔV^{cont} is the expected change in continuation value for the agent if he decides to not be truthful. The resulting distortions are illustrated in Figure 12 for the case of $\alpha = 1$. First, discrimination against the best proposals lowers the payoff to lying and eliminates any incentive to exaggerate for the highest-quality projects. General favoritism increases the continuation value of the agent. Finally, focused leniency is used to push the reneging temptation down to the continuation value for the region where the continuation value alone is not enough to keep truth-telling incentive-compatible.

In the case of uniform distribution and $\alpha = 1$, we can write the payoffs of the agent and the

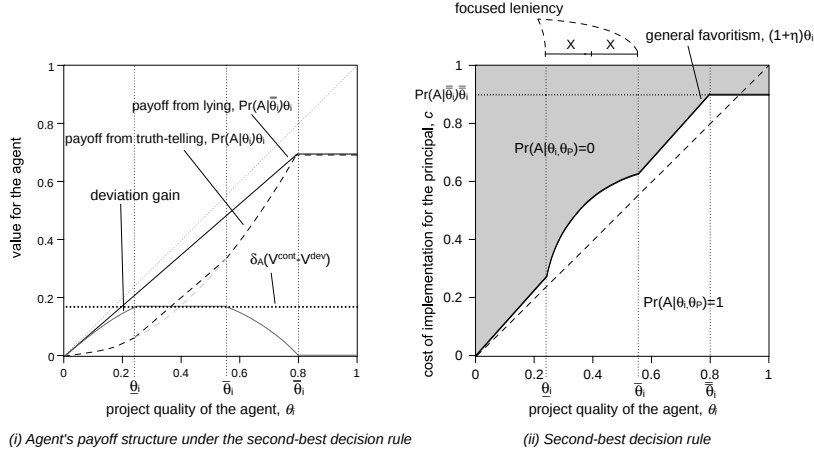


Figure 12: Example of a second-best decision rule

principal in a closed form as a function of only η and $\bar{p}$, the level of general favoritism and the maximal acceptance probability.¹⁰ The resulting continuation value then pins down the extent of focused leniency that is needed to maintain truth-telling by the agent.

The solution is illustrated in Figure 13. The key is panel (ii), which illustrates how the distortions grow as the agent becomes increasingly impatient, where both the range $[\underline{\theta}, \bar{\theta}]$ over which the principal chooses to exercise focused leniency and the range $[\bar{\theta}, 1]$ that the principal discriminates against are growing in the agent's impatience, worsening the equilibrium performance. In short, the increasing impatience of the agent leads the principal to shift the agent's influence from high-quality to average-quality projects. Relatedly, panel (iii) illustrates the extent of general favoritism, where the principal first increases the degree of favoritism to increase the continuation value for the agent, but once the agent becomes sufficiently impatient, the level of general favoritism is decreased because the low value that the agent places on the future makes the higher continuation value an inefficient means of providing value.

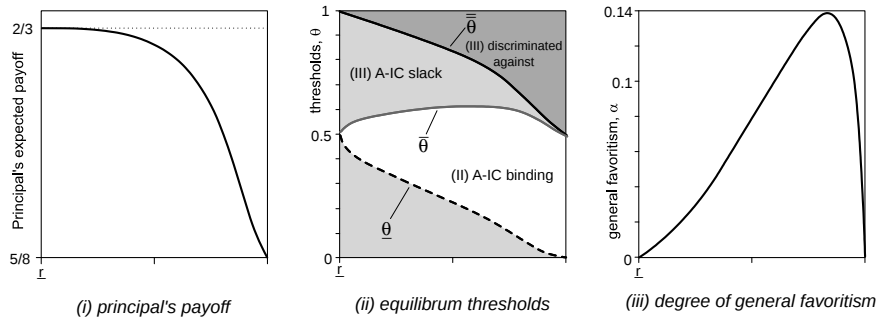


Figure 13: Principal's expected payoff and equilibrium distortions under commitment

The distortions are even more clearly highlighted if we consider how the actual equilibrium

¹⁰The derivation of these payoffs is not particularly instructive and available from the author on request.

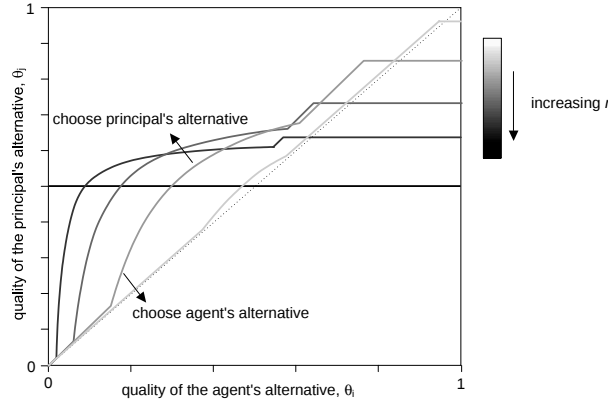


Figure 14: Equilibrium relational influence under principal commitment

decision rule changes as we change the agent's patience. These decision rules are illustrated in Figure 14, which plots the optimal decision rule of the principal for various discount rates of the agent. Intuitively, as the agent initially becomes more impatient, the principal uses all three tools at his disposal. She decreases the maximum acceptance probability to decrease the incentives to exaggerate, thus increasing the discrimination against the best alternatives of the agent. At the same time, to limit the rate at which such discrimination needs to grow, the principal increases the leniency towards the average proposals, increasing the "bulge" in the middle, in relation to the maximum acceptance probability. Finally, the principal initially increases the degree of general favoritism to further increase the agent's continuation value, but eventually decreases it when continuation value becomes a secondary concern to the agent. As the agent becomes infinitely impatient, the decision rule converges to the static optimum: the principal accepts any proposal by the agent whenever her own alternative is worse than average, while choosing her own proposal otherwise.

B.3 Proofs and derivations for the continuous state

B.3.1 Proof of Lemma 14

To find the least-cost means of increasing the continuation value, recall that the agent's expected utility is given by

$$\int_0^1 (b(1 - \alpha) + \alpha\theta_i) \Pr(A|\theta_i) f_A(\theta_i) d\theta_i,$$

while the cost of the distortions to the principal (here focusing on favoritism) is given by

$$\int_0^1 \left[\int_0^1 (\theta_i - c) \Pr(A|\theta_i, \theta_P) f_P(c) dc \right] f_A(\theta_i) d\theta_i$$

Two observations follow. First, the minimum cost of delivering any particular $\Pr(A|\theta_i)$ is for the principal to implement the project as long as $c \leq \Pr(A|\theta_i)$ while rejecting the project otherwise. Thus, we can define the expected probability of acceptance by the threshold rule $\bar{c}(\theta_i)$.

Second, on the margin of $\bar{c}(\theta_i)$, consider increasing $\bar{c}(\theta_i)$ while decreasing $\bar{c}(\theta'_i)$ in a way that keeps the expected cost to the principal constant. This implies that

$$(\theta_i - \bar{c}(\theta_i)) f_P(\bar{c}(\theta_i)) f(\theta_i) + (\theta'_i - \bar{c}(\theta'_i)) f(\theta'_i) f_P(\bar{c}(\theta'_i)) \frac{d\bar{c}(\theta'_i)}{d\bar{c}(\theta_i)} = 0 \Leftrightarrow \frac{d\bar{c}(\theta'_i)}{d\bar{c}(\theta_i)} = -\frac{(\theta_i - \bar{c}(\theta_i)) f(\theta_i) f_P(\bar{c}(\theta_i))}{(\theta'_i - \bar{c}(\theta'_i)) f(\theta'_i) f_P(\bar{c}(\theta'_i))}.$$

The corresponding effect on the agent's expected payoff is given by

$$\int_0^1 \left[\int_0^{\bar{c}(\theta_i)} (b(1-\alpha) + \alpha\theta_i) f_P(c) dc \right] f_A(\theta_i) d\theta_i$$

$$(b(1-\alpha) + \alpha\theta_i) f_P(\bar{c}(\theta_i)) f_A(\theta_i) + (b(1-\alpha) + \alpha\theta'_i) f_P(\bar{c}(\theta'_i)) f_A(\theta'_i) \frac{d\bar{c}(\theta'_i)}{d\bar{c}(\theta_i)}.$$

Now, holding the principal's payoff constant, $\frac{d\bar{c}(\theta'_i)}{d\bar{c}(\theta_i)} = -\frac{(\theta_i - \bar{c}(\theta_i)) f_A(\theta_i) f_P(\bar{c}(\theta_i))}{(\theta'_i - \bar{c}(\theta'_i)) f_A(\theta'_i) f_P(\bar{c}(\theta'_i))}$, and so optimality of the policy requires that

$$(b(1-\alpha) + \alpha\theta_i) f_P(\bar{c}(\theta_i)) f_A(\theta_i) - (b(1-\alpha) + \alpha\theta'_i) f_P(\bar{c}(\theta'_i)) f_A(\theta'_i) \frac{(\theta_i - \bar{c}(\theta_i)) f_A(\theta_i) f_P(\bar{c}(\theta_i))}{(\theta'_i - \bar{c}(\theta'_i)) f_A(\theta'_i) f_P(\bar{c}(\theta'_i))} = 0.$$

This then rearranges then to

$$\bar{c}(\theta'_i) = \bar{c}(\theta_i) + (\theta'_i - \theta_i) \frac{[b(1-\alpha) + \alpha\bar{c}(\theta_i)]}{(b(1-\alpha) + \alpha\theta_i)}.$$

Alternatively, we can benchmark this to zero, which gives us

$$\bar{c}(\theta_i) = \bar{c}(0) + \theta_i \left(1 + \frac{[\alpha\bar{c}(0)]}{b(1-\alpha)} \right),$$

which then allows us to write the relationship between any two states also as

$$\frac{\bar{c}(\theta'_i) - \bar{c}(\theta_i)}{\theta'_i - \theta_i} = \left(1 + \frac{[\alpha\bar{c}(0)]}{b(1-\alpha)} \right) > 1. \quad (47)$$

As $\alpha \rightarrow 1$, the problem becomes ill-defined, and we can go above to note that the expression becomes $\frac{\bar{c}(\theta'_i)}{\bar{c}(\theta_i)} = \frac{\theta'_i}{\theta_i}$, so the relationship is satisfied for $\bar{c}(\theta_i) = (1 + \eta)\theta_i$.

B.3.2 Leniency towards average proposals

Given the level of general favoritism, $\bar{c}(\theta_i|\alpha, b)$ and the degree of discrimination against the best proposals, $\bar{p}(\alpha, b)$, we can write the truth-telling constraint (without any additional modifications) as

$$\delta_A \Delta V^{cont} \geq \left(1 - \frac{F_P(\bar{c}(\theta_i|\alpha, b))}{\bar{p}(\alpha, b)} \right) u_A(\theta_i). \quad (48)$$

Now, let the cost distribution be uniform. Then, we can write the above as

$$\delta_A \bar{p}(\alpha, b) \Delta V^{cont} \geq (\bar{p}(\alpha, b) - \bar{c}(\theta_i|\alpha, b)) (b(1 - \alpha) + \alpha\theta_i). \quad (49)$$

First, finding the maximal deviation temptation, maximizing the right-hand side gives us

$$(\bar{p}(\alpha, b) - \bar{c}(\theta_i|\alpha, b)) \alpha - \left(1 + \frac{[\alpha \bar{c}(0)]}{b(1 - \alpha)}\right) (b(1 - \alpha) + \alpha\theta_i) = 0, \quad (50)$$

which rearranges to

$$\hat{\theta}_i = \frac{1}{2} \left(\frac{(\bar{p}(\alpha, b) - \bar{c}(0))}{\left(1 + \frac{\alpha \bar{c}(0)}{b(1 - \alpha)}\right)} - \frac{b(1 - \alpha)}{\alpha} \right). \quad (51)$$

I am making this intermediate observation by noting how the changes in the principal's decision rule already alter both the state for which the deviation temptation is maximized and the resulting gain to deviating. In particular, decreasing the maximum acceptance probability, $\bar{p}(\alpha, b)$, increasing the baseline favoritism, $\bar{c}(0)$, and the rate at which the favoritism grows, $\left(1 + \frac{\alpha \bar{c}(0)}{b(1 - \alpha)}\right)$, all both decrease the project quality for which the deviation temptation is maximized and the gain from deviating in the first place.

Now, if the condition is still not satisfied, we can use the idea of focused leniency. To this end, note that if the agent's IC constraint is not satisfied for $\hat{\theta}_i$, as long as the solution is interior, we have that there exists $\underline{\theta}_i < \hat{\theta}_i < \bar{\theta}_i$ for which the condition is exactly satisfied. These thresholds are given by the solution to

$$\begin{aligned} & [\delta_A \bar{p}(\alpha, b) \Delta V^{cont} - (\bar{p}(\alpha, b) - \bar{c}(0)) b(1 - \alpha)] \\ & - \theta_i [(\bar{p}(\alpha, b) - 2\bar{c}(0)) \alpha - b(1 - \alpha)] + \left(1 + \frac{[\alpha \bar{c}(0)]}{b(1 - \alpha)}\right) \alpha \theta_i^2 = 0. \end{aligned}$$

Given the threshold, the interior must then satisfy (since the deviation loss is the same)

$$\Pr(A|\theta_i) = \bar{p}(\alpha, b) - (\bar{p}(\alpha, b) - \bar{c}(\underline{\theta}_i|\alpha, b)) \frac{u(\underline{\theta}_i)}{u(\theta_i)} > \bar{c}(\theta_i|\alpha, b). \quad (52)$$

When the agent cares only about the quality of the project, the bounds simplify to

$$\{\underline{\theta}_i, \bar{\theta}_i\} = \frac{\bar{p}}{2(1 + \eta)} \left(1 \pm \sqrt{1 - 4 \frac{\delta_A(1 + \eta)}{\bar{p}} \Delta V^{cont}} \right). \quad (53)$$