

Essays on Repeated Principal-Agent Interaction

by

Shahin Baghirov

A Dissertation Submitted to the
Graduate School of Social Sciences and Humanities
in Partial Fulfillment of the Requirements for
the Degree of

Doctor of Philosophy

in

Economics



KOÇ ÜNİVERSİTESİ

July 6, 2022

Essays on Repeated Principal-Agent Interaction

Koç University

Graduate School of Social Sciences and Humanities

This is to certify that I have examined this copy of a doctoral dissertation by

Shahin Baghirov

and have found that it is complete and satisfactory in all respects,
and that any and all revisions required by the final
examining committee have been made.

Committee Members:

Prof. Levent Kockesen (Advisor)

Assoc. Prof. Ayça Ebru Giritligil (Advisor)

Prof. Alp Atakan

Prof. Mehmet Yiğit Gürdal

Prof. Seda Ertay

Assoc. Prof. Özgür Yılmaz

Phd. Umut Keskin

Date:

ABSTRACT

Essays on Repeated Principal-Agent Interaction

Shahin Baghirov

Doctor of Philosophy in Economics

July 6, 2022

In a repeated interaction, one tries to maximize his/her current payoffs and consider the decision's future implications. Reputation concerns arise when there is incomplete information regarding the preferences of the counterparts. I theoretically and experimentally analyze the finitely repeated interaction between an agent and a principal. The principal can be thought of as a state official who holds either a one-period or a long-term project. The agent is a private firm that holds the necessary information for the project.

The principal holds a project which yields a payoff depending on the realized state of the world. He is not informed about the state of the world and hires an informed agent to delegate the authority. The agent may be a good type who shares the same preferences with the principal. Moreover, there is a positive probability that the agent may be a “bad type” who has misaligned preferences with the principal. The principal does not know the agent's type but knows the prior probability that the agent is a bad type. The agent is a long-term player, and the principal is either a short-term or a long-term player. In the short-term-principal setup, the long-term agent meets a new principal in each period of the interaction. If the principal is also a long-term player, then the agent and principal meet once and interact as long as the principal continues to hire the agent. In both setups, the principal observes the past choices of the agent when making the hiring decision. Hence, both types of agents have a long-term incentive to keep a high reputation level in addition to the short-term incentive to play the stage-game optimal action. The good agent wants to separate from the bad agent, whereas the bad agent wants to mimic the good type.

This thesis includes three chapters. The first chapter defines a benchmark model. The short-term principal does not hire the agent in the initial periods where there

are reputation incentives. The more periods are left to be played, the more the agent's reputation incentives. Hence, the agent tends to play the action that will increase her reputation level even if that action is not optimal in the stage game. The principal expects a negative payoff in such periods. The reputation incentives lead to surplus loss as the principal does not hire the agent in such periods. Hence, the "bad reputation" behavior is observed [Ely and Välimäki, 2003]. When the principal is also a long-term player, the harmful effect of reputation incentives can be avoided. The long-term principal hires the agent in periods with reputation concerns. He can compensate for the losses of reputation-building behavior in the future periods. Moreover, the long-term principal hires the agent for lower reputation levels. Hence, it is seen that a long-term interaction between the agent and the principal decreases the harmful effects of reputation-building behavior and improves equilibrium payoffs.

In the second chapter, I introduce endogenous stakes. Either the agent or the principal sets the allocation of the relative stakes. The agent prefers to start the career path with larger stakes, decreasing the stakes gradually. By starting large, the (good) agent creates an environment where future interaction becomes less valuable. Hence the incentive to take the right action in the current period prevails over reputation concerns. On the other hand, the principal allocates the stakes so that the interaction starts small. By doing so, the principal benefits from the reputation concerns of the agent. The principal updates his belief on the agent through interaction - stakes increase as the agent's reputation rises. Likewise, the social planner - whose objective is to maximize social welfare - prefers to start the interaction small. The findings highlight the optimal design of a career path under different conditions.

The third chapter theoretically and experimentally analyzes a slightly different repeated principal-agent game with varying relative stakes. I focus on the setup where there is a high probability that the principal and the agent have misaligned preferences (the agent is the bad type with sufficiently high probability). Repeated play becomes valuable in such a setting by improving the equilibrium payoff of the principal. The agent has reputation incentives that motivate her to take action matching the true state in the initial periods rather than maximize her period payoff. As Morris [2001] calls, the "discipline effect" of the reputation incentives benefits the principal. We show it is optimal for the principal to start the interaction small and increase the stakes gradually. The agent's reputation incentives are managed so that the reputation evolves slowly. I test these predictions in four treatments

via online experiments. Each period receives an equal stake in the first treatment. The interaction starts small in the other three treatments, and the stakes increase at different speeds. We show that the smaller the interaction starts, the higher the reputation incentives of the agent. More importantly, I show that the principal earns a higher payoff in starting-small (gradualism) treatments than the equal-stakes treatment. I contribute to the literature on gradualism by showing that it is a valuable tool for improving equilibrium payoffs in the principal-agent framework with asymmetric information.



ÖZETÇE

Tekrarlanan Asil-Vekil Etkileşimi Üzerine Makaleler

Shahin Baghirov

Ekonomi, Doktora

6 Temmuz 2022

Tekrarlanan bir etkileşimde katılımcılar, mevcut getirileri maksimize etmenin yanı sıra kararlarının gelecekteki etkilerini de dikkate alırlar. Muhatapların tercihleriyle ilgili eksik bilgi olduğunda reputasyon kaygıları ortaya çıkar. Bu tez, sonlu tekrarlanan asil-vekil ilişkisini teorik ve deneysel olarak analiz etmektedir. Asil, bir dönemlik veya uzun dönemli bir projeyi elinde bulunduran bir devlet görevlisi olarak düşünülebilir. Vekil ise proje için gerekli bilgilere sahip özel kurumdur.

Asil, olası getirileri gerçekleştiren duruma bağlı bir projeyi uygulayacaktır ve gerçek durumla ilgili bilgiye sahip değildir. O, projeyi devretmek için bilgili bir vekili işe alıp almama kararı verir. Vekil, asille aynı tercihleri paylaşan iyi tip olabilir. Ayrıca vekil, pozitif olasılıkla asil ile uyumsuz tercihleri olan kötü tiptir. Asil, vekilin tipini bilmemekte ancak vekilin kötü tip olma olasılığını bilmektedir. Vekil, uzun vadeli; asil ise kısa vadeli veya uzun vadeli bir oyuncudur. Kısa vadeli asil tasarımı her dönemde uzun vadeli vekil yeni bir asille etkileşime girer. Uzun vadeli asil tasarımı ise, asil vekili işe almaya devam ettiği sürece etkileşim devam eder. Her iki tasarımda da asil, işe alma kararını verirken vekilin geçmiş seçimlerini gözlemler. Vekilin periyot-optimal aksiyonu oynamak için kısa vadeli teşvike ek olarak reputasyonunu artırmak için uzun vadeli kaygıları vardır. İyi vekil, kötü vekilden farklılaşmak isterken, kötü vekil iyi tipi taklit etmek istemektedir.

Bu tez, üç bölümden oluşmaktadır. İlk bölüm, bir ana model tanımlar. Kısa vadeli asil tasarımı vekil, reputasyon kaygılarının olduğu ilk dönemlerde işe alınmaz. Oynanacak ne kadar çok periyot varsa, vekilin o kadar fazla reputasyon kaygısı vardır. Bu nedenle, vekil, belirli periyotta optimal olmayan eylemi reputasyonunu artırmak için oynayabilir. Reputasyon kaygıları, asilin bu periyotlarda vekili işe almamasına yol açabilir. Bu nedenle, “kötü reputasyon” davranışı gözlemlenir [Ely and Välimäki, 2003]. Asil uzun vadeli bir oyuncu olduğunda, rep-

utasyon kaygılarının faydalı olduđu gözlemlenmektedir. Uzun vadeli asil, reputasyon kaygısı olan dönemlerde de vekili işe alabilir. Ayrıca, uzun vadeli asil, vekili daha düşük reputasyon seviyelerinde de işe almaktadır. Böylece, denge getirileri iyileşmektedir.

İkinci bölümde endojen ağırlıklar tanıtılmaktadır. Göreli ağırlıkları vekil, asil veya sosyal planlamacı belirler. Vekil, kariyer yoluna daha büyük adımlarla başlamayı tercih eder ve ağırlıkları kademeli olarak azaltır. İyi tip vekil, büyük başlayarak, gelecekteki etkileşimin daha az değerli hale geldiği bir ortam yaratır. Dolayısıyla, şimdiki dönemde doğru eylemi yapma teşviğı, reputasyon endişelerinin önüne geçer. Öte yandan, asil, etkileşimin küçük başlamasını tercih eder. Bunu yaparak asil, vekilin reputasyon kaygılarından yararlanır. Asil, etkileşim ilerledikçe vekil ile ilgili inancını günceller. Aynı şekilde, amacı sosyal refahı maksimize etmek olan sosyal planlayıcı, etkileşime küçük başlamayı tercih eder. Bulgular, farklı koşullar altında bir kariyer yolunun optimal tasarımını vurgulamaktadır.

Üçüncü bölüm, değişen göreli ağırlıkla tekrarlanan bir asil-vekil oyununu teorik ve deneysel olarak analiz etmektedir. Asıl ve vekilin farklı tercihlere sahip olma olasılığının yüksek olduğu tasarıma odaklanıyorum (vekil, yeterince yüksek olasılıkla kötü tiptir). Tekrarlanan oyun, böyle bir ortamda, denge getirisini geliştirerek değerli hale gelir. Asil için etkileşime küçük başlamanın ve ağırlıkları kademeli olarak artırmanın optimal olduğunu gösteriyorum. Vekilin reputasyon teşvikleri, reputasyonun yavaş gelişeceği şekilde yönetilir. Bu tahminleri çevrimiçi deneylerle dört farklı tretmanda test ediyorum. İlk tretmanda, her periyot eşit ağırlık almaktadır. Etkileşim, diğer üç tretmanda küçük başlar ve ağırlıklar farklı hızlarda artar. Etkileşim ne kadar küçük başlarsa, vekilin reputasyon kaygılarının o kadar yüksek olduğunu gösteriyorum. Daha da önemlisi, asil küçük başlanılan tretmanlarda eşit ağırlıklı tretmana göre daha yüksek bir getiri elde etmektedir. Asimetrik bilgi içeren asil-vekil iletişiminde denge getirilerini geliştirmek için “küçük başlamanın” değerli bir araç olduğunu göstererek literatüre katkıda bulunuyorum.

ACKNOWLEDGMENTS

I have received a great deal of support and assistance throughout the writing of this dissertation.

First and foremost, I would like to express my indebtedness and a special thanks to Professor Levent Koçkesen whose guidance and expertise were priceless throughout my Ph.D. study. I benefited from his passion, dedication, and, most notably, his support which pushed me to sharpen my thinking and brought my work to a higher level. I feel immensely privileged to have worked so closely with him.

I also want to say special thanks to Associate Professor Ayça Ebru Giritligil. She did not only help me to bring a different insight into my research but also provided invaluable support for me, making my path to Ph.D. much easier. I consider myself extremely lucky to have worked with her and her BELIS team.

I would like to acknowledge Professor Alp Atakan, Professor Seda Ertaç, and Associate Professor Mehmet Yiğit Gürdal for their kind interest in my research and helpful guidance throughout my thesis progress. Their suggestions and feedback helped me a lot to enhance my research.

I extend my genuine thanks to all members of the Department of Economics at Koç University and Istanbul Bilgi University who contributed directly or indirectly to this thesis. I would also like to thank participants at seminars and conferences that I have attended worldwide for valuable recommendations and remarks.

I would like to thank my mother, Filura, for her wise advisory and understanding ear. Without her full support, I would not find the courage to start and pursue an academic career. I owe a heartfelt thanks to my precious sisters Şahane and Meltem for their love and faith during all these years. Most importantly, I want to thank my beloved wife, Burcu, for not stopping to believe in me for a moment.

TABLE OF CONTENTS

List of Tables		xii
List of Figures		xiii
Chapter 1:	Repeated Principal-Agent Interaction	1
1.1	Introduction	1
1.2	Literature	3
1.3	The Model	5
1.3.1	One Period Model	11
1.4	Short-term Principal	12
1.5	Long-term principal	19
1.5.1	What happens when n increases?	23
1.5.2	What happens when k increases?	24
1.5.3	When there will be reputation concerns?	25
1.5.4	Commitment bad agent	27
1.6	Conclusion	30
Chapter 2:	Endogenous Stakes in a Repeated Principal-Agent Interaction	32
2.1	Introduction	32
2.2	Literature	35
2.3	The Model	36
2.4	Short-term Principal	39
2.4.1	Agent sets stakes	39
2.4.2	Social planner sets stakes	41

2.5	Long-term principal	45
2.5.1	Agent sets stakes	45
2.5.2	The long-term principal sets relative stakes	51
2.6	Conclusion	57
Chapter 3: Gradualism in a Principal Agent Interaction - Experiment		60
3.1	Introduction	60
3.2	Literature Review	62
3.3	The model	65
3.4	Experimental Design	67
3.5	Theory and Predictions	71
3.5.1	Reputation Building Behavior	71
3.5.2	Payoffs	77
3.6	Results	77
3.6.1	Agent Behavior and Reputation Incentives	78
3.6.2	Beliefs	79
3.6.3	Principal Behavior	81
3.6.4	Payoff Comparison	83
3.7	Conclusion	85
Appendix A: Chapter 1 Proofs		95
Appendix B: Chapter 2 Proofs		113
Appendix C: Chapter 3 Comparison of Treatment parameters		136
Appendix D: Chapter 3 Instructions and Screenshots (Treatment 4)		141
D.1	English Instructions	141
D.2	Turkish Instructions and Screenshots	148



LIST OF TABLES

1.1	Payoff table.	7
1.2	Strategies table.	8
3.1	Payoff table.	66
3.2	Strategies table.	66
3.3	Reputation incentives across treatments.	87
3.4	Belief change.	88
3.5	Belief update.	89
3.6	Frequency of subjective believes in the state.	89
3.7	Hiring behavior.	90
3.8	Payoff comparison.	90
C.1	Treatment 1 payoff table.	136
C.2	Treatment 2 payoff table.	137
C.3	Treatment 3 payoff table.	138
C.4	Treatment 4 payoff table.	139
C.5	Treatment parameters.	140

LIST OF FIGURES

2.1	The equilibrium allocation of stakes when the agent set stakes. $\underline{\delta} = 0.1$ and $n = 100$	42
2.2	The equilibrium allocation of stakes when the long-term principal set stakes and the bad agent is a commitment type. $\underline{\delta} = 0.1$ and $n = 100$	54
2.3	The equilibrium allocation of stakes when the long-term principal set stakes and the bad agent is a strategic type. $k = \frac{1}{4}$ and $n = 100$	57
3.1	The stake allocation for four treatments.	68
3.2	Change in believes. (01 corresponds to the observation of state 0 and action 1.)	81
3.3	Second period believes. (01 corresponds to the observation of state 0 and action 1.)	82

Chapter 1

REPEATED PRINCIPAL-AGENT INTERACTION

1.1 Introduction

I propose an uninformed principal's interaction with an informed and potentially bad agent for a finite number of periods. The principal can be described as an employer who does not have the necessary information. He needs an informed agent to delegate the authority. The principal might seek an employee for a one-time job or a project that will last for a finite number of periods. In each period of the relationship, a state of the world is realized. While the good agent favors the action that matches the state of the world, the bad agent prefers the higher action in the stage game.

In an influential paper, Ely and Välimäki [2003] (hereinafter EV) showed that reputational concerns might sometimes hurt an agent, a situation that they call “bad reputation”. More precisely, they analyzed an interaction between a short-term principal and a long-term agent in an infinitely-repeated setup. They showed that if there is a positive probability of the agent being a bad type, then the principal does not hire the agent in the equilibrium. The reasoning is as follows. The good agent has a short-term incentive to take the correct action and a long-term incentive to increase her reputation. There are infinitely many periods to be played, hence the reputational incentives of the agent triumph. If there is a positive probability that the agent is a bad type, then the agent prefers to reveal herself as good to future employers. Thus, she chooses the action to increase her reputation rather than the necessary action. The principal knowing the reputation incentive of the agent, does not hire her. Therefore, the agent is not hired in all renegotiation-proof

Nash equilibria, leading to the loss of all surplus.

Infinite periods analysis seems more suitable in interaction with no predefined number of periods. However, in numerous real-world applications, the length of a project or career is somewhat known. The project's lifespan is known, especially when a long-term principal searches for an agent to delegate the authority over a project.

In this chapter, I analyze a finite period version of EV. I show that their results continue to hold in finite periods. More specifically, the agent is not hired in the initial periods. The reasoning is that if there are many periods to be played in the future, then the agent has high reputational incentives. Unlike EV, I have a finitely-repeated setup; hence, there is always a period after which no reputational incentives are left. If the number of periods tends to infinity, then the agent is never hired, as showed by EV.

The short-term principal does not hire the agent if she has reputation concerns in the period. The reputation concern of the agent depends on several parameters. Firstly, if the interaction length is sufficiently short, then the agent has no reputation incentive. It is not rational to take the reputation-driven action if not many periods are left to be played. Second, the value of the project has a direct effect on reputation concerns. The higher the potential profit from a project is, the more reputational concerns arise. Lastly, the agent's reputation level matters. An agent with a sufficiently high reputation level does not care about reputation-building much.

However, when the principal is also a long-term player, the principal hires the agent when the agent has reputation incentives. In addition, he hires the agent for lower initial reputation levels as expected. The principal gets a negative payoff in the periods where the agent takes the costly reputation-driven action. However, the principal benefits from the increase in the agent's reputation in the long term. Hence, he can compensate for the losses stemming from the agent's reputation-driven behavior. Thus I show that reputation concerns are not particularly bad when the principal is also a long-term player. EV also showed that we could reduce the harm of the reputation concerns by making the principal a long-term player.

1.2 Literature

A vast literature exists on the role of reputation in two-agent or principal-agent problems. Papers by Kreps and Wilson [1982] and Milgrom and Roberts [1982a] were the first to formalize the effects of the reputation that pioneered the conventional wisdom that the reputation concern increases commitment power. These papers analyzed models in which a long-term player meets with several short-term agents, and they showed that adding even some imperfect information leads to the rise in the reputation effect. An incumbent monopolist faces short-term entrants. If there is a small probability that the incumbent monopolist is tough, he may want to maintain a reputation for toughness to defer the other agents' entry. Kreps and Wilson [1982] and Milgrom and Roberts [1982a] showed that the introduction of the repeated actions and asymmetry in information leads to a reputation effect. In that framework, even if the predatory action is irrational for the monopolist in the stage game, he takes that action to gain a reputation as a predator. This action helps the monopolist to prevent short-term players from entering the market. He enjoys his monopolistic power in the future by investing in his reputation today.

Fudenberg and Levine [1989, 1992] applied these results to a more general set of games where a long-run player faces short-run opponents. Schmidt [1993] and Celetani et al. [1996] analyzed the role of reputation when there are two long-run players. Schmidt concludes that if the first player is sufficiently more patient than the second player, the reputation benefits the first player. Likewise, Celetani et al. [1996] showed that if a player (whose type is not known by the other player) is more patient than his opponent, he can exploit the less patient player's uncertainty by holding a reputation as a commitment type. This literature following Kreps and Wilson [1982] and Milgrom and Roberts [1982a] suggested that reputation building is an effective tool to increase payoff compared to an environment with no reputation incentives.

Holmström [1999] defined the reputation incentive as an employee's career motives and showed that the reputation concerns could be beneficial or harmful depend-

ing on the level of conflict (or alignment) between the employer and the employee. Morris [2001] analyzed a twice repeated interaction in a cheap-talk environment with the noise. He defines three effects of the reputation incentives. Firstly, the “discipline effect” of the reputation concerns leads the bad agent to choose an action that hurts the principal less. In our setup, this effect leads to the bad agent playing the action that matches the state of the world, rather than always playing the higher action. Secondly, through the “sorting effect”, the principal can update his belief about the agent through the relation. The first two effects of the reputation concerns improve the equilibrium payoffs. However, the third effect, “political correctness” leads the good agent to choose the action that will increase her reputation rather than choosing the right action. All in all, Morris argues that the reputation concerns’ overall effect is unambiguous.

In the model proposed by Ely and Välimäki [2003], a short-term uninformed principal decides whether or not to hire a possibly bad informed agent defined as a long-term player. They represented two types of agents, rational and bad. The bad type was committed to taking the bad action regardless of the true state. They found that the good (rational) agent tended to take the good action to be revealed as good rather than taking the necessary action. By doing so, the agent increased his reputation so that he would be hired in the future. Thus, it led to losing the current principal if he hired the agent. EV showed that if there is a positive probability of the agent being a bad type, then the principal’s utility is bound by a value that converges to 0 if the discount factor is close enough to 1. They found that the reputation effect may lead to a loss of all surpluses.

Ely and Välimäki [2003] continued to show that the losses from a bad reputation could be avoided if the principal is also a long-term player. If the principal and the agent are long-term players, the agent goes through an evaluation phase. In the evaluation phase, the agent is hired irrespective of the action observed. If the reputation of the agent gets above some level, then she is hired for some periods. I show that the equilibrium behavior is similar in a finitely-repeated interaction.

Ely et al. [2009] defined a more general setting to understand when a game

is a “bad reputation” game. They had a richer type set in this setting. They showed that if the probability of an “unfriendly” commitment type is sufficiently high and the long-run player is patient enough, then the utility of the long-run player converges to exit payoff in any equilibrium. Then, the effect of the reputation is “bad”. Unfriendly commitment types are the agents who commit to bad actions. Moreover, Ely et al. [2009] showed that if a Stackelberg type is likely enough, then the bad reputation result does not hold.

1.3 The Model

This section will define the basic model where the stakes are equal across periods. I do not have a discounting factor. Ely and Välimäki [2003] provide the analysis with discounting close enough to 1. Hence, I assume a discounting factor equal to 1 to avoid future confusion when I introduce the stake variable in chapter 2.

An informed long-term agent meets an uninformed principal each period over n periods. The principal could be short-term or long-term. I start with the short-term principal. Let N be the index set of the periods. In each period $i \in N$, the state $\theta_i \in \{0, 1\}$ is realized. Each state is equally likely, i.e. $\text{prob}[\theta_i = 0] = \text{prob}[\theta_i = 1] = 1/2$ and independent across the periods. The agent observes the state, whereas the principal does not.

At the beginning of each period, the agent meets with a new principal. The principal decides whether to hire the agent or not. Period i principal’s action is defined by λ_i , where $\lambda_i = 0$ means he does not hire the agent and $\lambda_i = 1$ means the principal hires her. If the principal does not hire the agent, the agent moves to the next period and meets with another principal. If $\lambda_i = 1$, then the agent observes the true state and takes action $a_i \in \{0, 1\}$. The payoffs are realized, and the agent moves to the next period.

There are two types of agents: $\{g, b\}$ - *good* and *bad* types. The good agent shares the same preferences with the principal, while the bad agent favors action 1. The agent is privately informed about his type. $\rho_1 \in (0, 1)$ measures the initial reputation of the agent being good.

For any $i \in N$, let H_i be the set of all histories before decision i is made, i.e. $H_i = (o_1, o_2, \dots, o_{i-1})$ where $o_i = (\lambda_i, \alpha_i)$. $\alpha_i \in \{\emptyset, \{0, 1\}\}$ denotes the outcome of the period i action. If the agent is not hired in period i , then $\alpha_i = \emptyset$. The period i strategy of the principal is given by $\Lambda_i : H_i \rightarrow \{0, 1\}$ where $\Lambda_i(h)$ is the principal's choice of λ_i . Principal's belief about the type of agent is defined by the probability of agent being good in period i : $\rho_i : H_i \rightarrow [0, 1]$. The principal observes all the previous actions of the agent but does not observe the true state in these periods. The good agent observes the state of world before taking action but the bad agent does not. The good and bad agent moves after histories of this type: $I_i^g = (h, \lambda_i, \theta_i)$ and $I_i^b = (h, \lambda_i)$, respectively, where I_i^a is the period i information set of the agent type $a \in \{g, b\}$. The good and the bad agent's (mixed) strategy is defined by the probability of agent playing action 0, given by $\mu_i(h, 1, \theta_i)$ and $\nu_i(h, 1)$, respectively. $\mu_i \& \nu_i : \mathfrak{I}_i \rightarrow \Delta[0, 1]$ where $\mathfrak{I}_i$ is the set of all period i information sets.

The timing of the stage game in period i can be summarized as follows:

1. The principal makes the hiring decision. If $\lambda_i = 0$, then the game moves to the next period. Both parties get a payoff of 0.
2. If $\lambda_i = 1$, nature chooses the state $\theta_i \in \{0, 1\}$.
3. The agent observes θ_i and takes the action $a_i \in \{0, 1\}$

The stage game payoffs are given by the Table 1.1 where $k \in (0, 1/2)$ is assumed to be the payoff accruing from the project in each period. This constraint on k is made for the following reason. In the one-shot game where there are no reputational incentives, the principal's stage game payoff is given by $\rho k + (1 - \rho)(k - 1/2)$. If the principal believes that $\rho = 0$, i.e., the agent chooses her action independent of the state, he will not hire the agent. If the agent is not hired, all the players get 0 payoffs. If the payoff received from the project (k) is "not sufficiently high", then the outside option is a better choice for the principal than facing an agent whose action does not depend on the state. In other words, the agent is always hired if $k \geq 1/2$.

Table 1.1: Payoff table.

Good agent (and principal)			Bad agent		
	$\theta_i = 0$	$\theta_i = 1$		$\theta_i = 0$	$\theta_i = 1$
$a = 0$	k	$k - 1$	$a = 0$	$k - 1$	$k - 1$
$a = 1$	$k - 1$	k	$a = 1$	k	k

The expected continuation payoff of the agent type $A \in \{g, b\}$ (good and bad agent, respectively) in period i is given as follows: $V_i^A = \sum_{j=1}^n U_j^A$ where U_i^A denotes the stage game payoff. As period i principal interacts with the agent for one period, he aims to maximize U_i^P .

Period i stage game strategy of the agents is summarized with the Table 1.2. $\mu_i(h, 1, \theta_i)$ and $\nu_i(h, 1)$ is denoted by $\mu_i^{\theta_i}$ and ν_i for the ease of exposition. $\mu_i^{\theta_i}$ denotes the probability with which the good agent plays action 0 given $\theta_i \in \{0, 1\}$ and ν_i denotes the probability with which the bad agent plays action 0 in any state of the world. I assume that the bad agent does not observe the state of the world. She is the type that mixes between two actions irrespective of the state of the world. I define equilibria for both commitment and strategic bad types. Commitment bad type agent always plays stage-game maximizing strategy, which corresponds to $a_i = 1 \forall i \in N$. Note that $\nu_i = 0 \forall i \in N$ if the bad agent is a commitment type. The strategic bad type is a payoff-maximizing player who does not necessarily play the stage-game optimal action.

Given the above-defined strategies, the principal's belief about the type of the agent in period $i + 1$ is updated by Bayes' rule as following:

$$\rho_{i+1}(a_i = 0|h) = \frac{\rho_i(\mu_i^0 + \mu_i^1)}{\rho_i(\mu_i^0 + \mu_i^1) + 2(1 - \rho_i)\nu_i}$$

$$\rho_{i+1}(a_i = 1|h) = \frac{\rho_i(2 - \mu_i^0 - \mu_i^1)}{\rho_i(2 - \mu_i^0 - \mu_i^1) + 2(1 - \rho_i)(1 - \nu_i)}$$

Table 1.2: Strategies table.

	$\theta_i = 0$	$\theta_i = 1$
Good agent	0 with probability μ_i^0	0 with probability μ_i^1
	1 with probability $1 - \mu_i^0$	1 with probability $1 - \mu_i^1$
Bad agent	0 with probability ν_i	
	1 with probability $1 - \nu_i$	

I focus on the perfect Bayesian equilibria with Markovian property. In that regard, history matters only in terms of its effect on the agent's reputation. Thus, if the agent has the same reputation level after two different histories, both the agent's and the principal's actions are the same in both cases. The decisions depend on the past observations only through their effect on the agent's reputation.

Markovian Property. For any $i \in N$ and $h, h' \in H_i$: $\rho_i(h) = \rho_i(h')$ implies $\lambda_i(a_i|h) = \lambda_i(a_i|h')$, $\mu_i(h) = \mu_i(h')$ and $\nu_i(h) = \nu_i(h')$.

For notational simplicity, I will suppress h in the notation and denote the action of the principal after observing a_i given h by $\lambda_i(a_i)$, i.e. $\lambda_i(a_i) \equiv \lambda_i(a_i|h)$. In the same manner, $\rho_{i+1}(x) \equiv \rho_{i+1}(a_i = x|h)$ where $x \in \{0, 1\}$.

I use agent-optimality as the equilibrium selection criteria when there are multiple equilibria. The agent-optimal equilibrium is the equilibrium that leads to the highest expected continuation payoff for the agent among all equilibria.

Define reputation concern as taking an action a_i other than stage-game payoff maximizing action which lead to $\rho_{i+1}(a_i) > \rho_i$. In other words, if the agent has a reputation concern in period i , then she has an incentive to play the action that will contribute to her reputation level, even if that action hurts the stage-game payoffs. Playing such an action with a positive probability is only sequentially rational if the agent gets a sufficiently high payoff following that action. I show that reputation-driven action is action 0.

First, note that the agent's payoff is increasing in reputation level in any agent-optimal equilibrium. This condition may seem natural, but there are equilibria in which a lower reputation level leads to a higher payoff to the agent. Consider period i such that no reputation concern in period i onward is an equilibrium. In other words, the good agent plays the action matching the state of the world, and the bad agent plays action 1. Now suppose there is an equilibrium in which for some ρ'_i , both types of agents play action 1 with probability 1 if $\rho_i > \rho'_i$. Both types of agents play the stage-game payoff maximizing action for lower reputation levels, i.e., no reputation concern. Given $\rho_i > \rho'_i$, define the principal's off-the-equilibrium path belief as $\rho_{i+1}(0) = 0$. This assessment may constitute an equilibrium. In that case, the agent may get a higher payoff if $\rho_i \leq \rho'_i$ (compared to $\rho_i > \rho'_i$). Hence, the agent does not always prefer a higher reputation level in *any* equilibrium. I show that a higher reputation level implies a higher payoff for the agent in the agent-optimal equilibrium.

Suppose for the contradiction that action 0 decreases the agent's reputation in period i . As the agent prefers a higher reputation level, the bad agent plays action 1 with probability 1 in such period. If good agent plays action 0 with a positive probability, then $\rho_{i+1}(0) = 1$ by Bayes's rule, contradicting $\rho_{i+1}(0) < \rho_i$. Hence, action 0 can only decrease the agent's reputation if it is played with 0 probability, i.e., not observed on the equilibrium path. There are such babbling equilibria where both types of agents play action 1 and the principal puts a sufficiently low belief on the probability of the agent being good type after action 0. No type of agent has an incentive to deviate to play action 0 if the future is sufficiently important. In such case, the agent takes the costly action 1 (when $\theta_i = 0$) and her reputation is not updated ($\rho_{i+1}(1) = \rho_i$). I show that there always exists a separating equilibrium which implies a higher payoff to the agent. The Lemmas 1.1, 1.2 and 1.3 formalize the argument.

Lemma 1.1. *In the agent-optimal perfect Bayesian equilibrium,*

$$\rho_{i+1}(a_i = 0|h) \geq \rho_i \geq \rho_{i+1}(a_i = 1|h) \quad \forall i \text{ and } \forall h \in H_i$$

which holds strictly when $\rho_i \neq 1$.

Proof. Appendix A. □

If $\theta_i = 0$, then the good agent will play action 0 with probability 1 ($\mu_i(h, 1, 0) = 1$). Note that the good agent has no incentive to act differently from the state observed when $\theta_i = 0$. Thus, playing action 1 with a positive probability when the state is 0 can not be an equilibrium strategy for the good agent. It is easy to see as playing action 1 hurts both the stage game payoff and the reputation. When the true state is 1, she has a reputational incentive to play 0 as it will increase her reputation as a good agent. Hence,

Lemma 1.2. *In the agent-optimal perfect Bayesian equilibrium, $\mu_i(h, 1, 0) = 1 \ \forall i$ and $\forall h \in H_i$.*

Proof. Appendix A. □

Lemma 1.1 and Lemma 1.2 imply that ν_i can not be above a certain level in any equilibrium. The reasoning is straightforward. If the bad agent plays action 0 with a sufficiently high probability compared to the good agent, then $\rho_{i+1}(1) > \rho_{i+1}(0)$ may be the case. But then, it implies that the agent has no motivation to play action 0. Thus, the agent will deviate to play action 1 with probability 1. Hence,

Lemma 1.3. *In the agent-optimal perfect Bayesian equilibrium, $\nu_i \leq \frac{1 + \mu_i}{2}$.*

Proof. Appendix A. □

Thus, I can interpret the strategies μ_i and ν_i as measuring the reputation concerns of the two types of agents. For ease of exposition, I will write μ_i to refer to $\mu_i(h, 1, 1)$ because $\mu_i(h, 1, 0) = 0$ in the agent-optimal equilibrium.

The reputation concerns depend on n . It corresponds to the length of agent's lifespan in the short-term principal setup and the length of the interaction in the long-term principal setup. The higher the number of periods is left to be played in the future, the (weakly) more reputation concerns the agent will have. EV analyzes

infinitely repeated setup; hence there are always reputation concerns if there is a positive probability that the agent is bad. However, in the finitely repeated setup, there is always a period i after which there is no more reputation concerns. Consider a critical period i after which an agent is hired only if $a_i = 0$ is observed. There are $n - i$ periods left in the future. Suppose $(n - i)k + k - 1 \geq k$ and $\mu_i = 0$ and $\nu_i = 0$. Note that $\rho_{i+1}(0) = 1$, and hence, both types of agents have an incentive to deviate to play action 0. Hence, $\mu_i = 0$ and $\nu_i = 0$ can not be equilibrium strategies. On the other hand, if $(n - i)k + k - 1 < k$, then $\mu_i = 0$ and $\nu_i = 0$ in any equilibrium. As playing action 0 leads a maximum of $(n - i)k + k - 1$, the agent prefers to play the stage-game maximizing action, which leads to a strictly higher payoff, k . Hence, $n < 1/k + i$ is the sufficient condition for the agent to not have reputation concerns in period i . The following proposition formalizes the argument for the initial reputation.

Proposition 1.1. *If $n < 1/k + 1$, then the agent has no reputational incentive in any period in any perfect Bayesian equilibria, i.e. $\mu_i \& \nu_i = 0 \forall i \in N$.*

Proof. Appendix A. □

Note that $k < 1/2$; hence, if the interaction lasts for three or fewer periods, then the agent will not have reputation concerns in any equilibrium. The same reasoning also holds for the last three periods of the interaction of any length.

1.3.1 One Period Model

I first provide the equilibrium analysis of the one-period model, i.e., $n = 1$. When the agent and the principal meet only once (e.g., the principal is a short-term player), the principal's objective is to maximize his stage game payoff. Thus, any reputational incentive of the good agent is harmful to the principal. Moreover, even with no reputational incentives, the principal does not hire the agent if the agent's reputation is below a certain level, more precisely $1 - 2k$. For the higher reputation levels, the principal hires the agent only if he believes the good agent plays action 0 with sufficiently low probability when the true state is 1. In other words, if the good

agent has a strict reputational incentive in a period, then she will not be hired. It follows from the following stage-game payoff of the principal:

$$U^P = \frac{1}{2}\rho k + \frac{1}{2}\rho\mu(k-1) + \frac{1}{2}\rho(1-\mu)k + (1-\rho)\left(k - \frac{1}{2}\right)$$

A strategic bad agent's strategy does not depend on the state. Hence, it does not affect the principal's payoff in a one-shot game. On the other hand, the principal's payoff decreases in μ . Given $\rho \geq 1 - 2k$, U^P is non-negative if μ is below μ' such that:

$$\mu' = \frac{2k + \rho - 1}{\rho}$$

Therefore, in a one-shot game, the agent is hired if and only if $\rho \geq 1 - 2k$ and $\mu \leq \mu'$. Regarding the agent's behavior in a one-shot game, it is clear that there will be no reputational incentives. Hence, $\mu \&\nu = 0$. Thus, the agent will be hired as long as the reputation is greater than or equal to $1 - 2k$.

1.4 Short-term Principal

This section applies the model proposed by Ely and Välimäki [2003] to a finite period framework. I start with the setup where the bad agent is a commitment type who plays action 1 in each period.

The behavior of the principal in a one-shot game is the same as the short-term principal. The following proposition formalizes the equilibrium behavior of the principal.

Proposition 1.2. *In any perfect Bayesian equilibrium:*

(a) *If $\rho_i < 1 - 2k$, then $\lambda_i = 0$.*

(b) *If $\rho_i \geq 1 - 2k$, then*

$$\lambda_i = \begin{cases} 0, & \mu_i > \mu'_i \\ 1, & \mu_i \leq \mu'_i \end{cases}$$

Proof. Appendix A. □

Define $P(t) = \frac{(1-2k)2^{t-2}}{(1-2k)2^{t-2} + k}$ where n denotes the total number of periods and $t \in \{1, 2, 3, \dots, n\}$. I will use $P(t)$ to determine the initial reputation's bounds when I characterize the equilibrium. If the agent's initial reputation is sufficiently high, then the agent has no reputation incentive in the unique equilibrium. In the equilibrium, the agent is hired as long as the reputation does not fall below $1-2k$. In particular, $\rho_1 \geq P(n)$ implies that if the agent never plays the costly reputational action ($\mu_i = 0$), then in the history where $\theta_i = 1$ and action 1 is observed for all periods, $\rho_n \geq 1-2k$. Hence, the agent is hired in all periods irrespective of history, including the last period.

Proposition 1.3. *If $\rho_1 \geq P(n)$, then the agent has no reputational incentive in the agent-optimal perfect Bayesian equilibrium. $\mu_i = 0$ and $\nu_i = 0 \forall i \in N$. The principal hires the agent for all periods: $\lambda_i = 1 \forall i \in N$.*

Proof. Appendix A. □

If, on the other extreme, $\rho_1 < P(1) = 1-2k$, the short-term principal never hires the agent even if the agent has no reputational incentives. Now I will describe the optimal equilibrium for the intermediate levels of the reputation.

First, I define necessary and sufficient conditions for the absence of reputation concerns. Irrespective of the initial reputation of the agent, if $n < \frac{1}{k} + 1$, then there are no reputation concerns in any equilibrium. Following Lemma 1.4 defines it for general initial reputation level.

Lemma 1.4. *$P(t) \leq \rho_1 < P(t+1)$ for $t \in \{1, 2, 3, \dots, n\}$. If $n < \frac{1}{k} + t$, then the agent has no reputational incentive in the agent-optimal perfect Bayesian equilibrium.*

$$\mu_i = 0 \text{ and } \nu_i = 0 \forall i \in N.$$

Proof. Appendix A. □

The short-term principal's payoff does not depend on the strategy of the bad type. The introduction of the strategic bad type does not directly affect the principal's

payoff. When making a hiring decision, the principal's only concern is μ_i and ρ_i . However, if the bad agent is also maximizing the expected continuation payoff, she will also play action 0 for some time. One may expect this to decrease the good agent's reputational incentives because playing action 0 contributes to the reputation level to a lesser degree. However, I can show that the equilibrium behavior does not change.

Only the reputational incentives change depending on whether the bad agent is strategic type or not. However, eventually, I get the same agent-optimal equilibrium with the commitment type. The reasoning is that the good agent has weakly higher reputational incentives than the bad agent in any period. In the model defined by EV, where the agent has a lifetime of infinity, the good agent has strictly higher reputation incentives. I will explain the difference later. First, I will provide the intuition for the higher reputational incentives of the good agent.

Lemma 1.5. *In the agent-optimal perfect Bayesian equilibrium, if the agent is hired in period i and $\nu_i > 0$, then the following must hold $\rho_{i+1}(0) \geq P(n-i) = \frac{(1-2k)2^{n-i-2}}{(1-2k)2^{n-i-2} + k}$. In words, if the bad agent mixes in period i , then the reputation after action 0 must be sufficient for the agent to be hired for all the remaining periods. Otherwise, the good agent has higher reputational incentives, hence $\nu_i > 0$ implies $\mu_i = 1$ and the agent is not hired.*

Proof. Appendix A. □

Consider a crucial period i in which the bad agent plays both actions with a positive probability. Suppose the agent survives for l periods after $a_i = 1$ and for m periods after $a_i = 0$. The bad agent is indifferent between two actions, hence, $mk - 1 = lk$ (which implies $m > l$). In each of the following periods, the bad agent may play action 1 with a probability of 1 or mix between the two actions. Consider the good agent. In each period, the state is 0 with $1/2$ probability. Therefore, the good agent may increase her reputation costlessly with a probability of $1/2$ in any period. As there are more periods to be played after $a_i = 0$ ($m > l$), bad agent

being indifferent ($mk - 1 = lk$) implies the good agent is indifferent between the two actions or strictly prefers action 0. There are two possible cases:

Case 1. $m < (n - i)$. It implies the agent does not survive until the end of the game if $a_i = 0$. In other words, another crucial period will be reached in the future. In this case, there will be reputational incentives in the periods following i . The future expected continuation payoff of the good agent exceeds that of the bad agent. In such case, $\nu_i > 0$ implies $\mu_i = 1$, and the agent is not hired in period i .

Case 2. $m = (n - i)$. If $a_i = 0$ is observed, the agent will be hired for all the remaining periods. Hence, there will be no reputational incentives in the remaining periods. Suppose the agent will be hired for all the periods after action 0 in the current period. In that case, the future expected continuation payoff of a good agent is equal to that of the bad agent. As the agent will be hired after any history, both types of agents will get $(n - i)k$ payoff. In this case, the good agent may also be indifferent between the two actions.

Let $P(t) \leq \rho_1 < P(t + 1)$ and suppose that the agent is first hired in period i where there is a reputational concern. Then $m = (n - i)$ must hold so that $\mu_i < 1$. The agent has no reputational incentives if first hired after the period $\lceil n - t + 1 - 1/k \rceil$. Hence, fix a period $i < n - t + 1 - \frac{1}{k}$ and suppose the agent survives for l periods after $a_i = 1$ is observed. The sequential rationality constraint of the bad agent implies $(n - i + 1)k - 1 = k + lk$. Simplifying the equality, I get $l = n - i - \frac{1}{k}$. $i < n - t + 1 - \frac{1}{k}$ implies $l > t - 1$. Therefore, $\rho_{i+1}(0)$ and $\rho_{i+1}(1)$ should be so that $m = n - i$ and $l = t - 1$. In words, if action 0 is observed in period i , then the agent is hired for the remainder of the game. On the other hand, if action 1 is observed, the agent is hired for $t - 1$ periods if no action 0 is observed. But, Bayes' rule states that given $P(t) \leq \rho_1 < P(t + 1)$, if action 1 is observed in each period, then the agent's reputation stays above $1 - 2k$ for at most t periods. Hence, l can be at most $t - 1$, i.e. $l \not> t - 1$. Hence, hiring does not start until the reputational incentives are eliminated. Thus, I define equilibrium where the bad agent is a strategic type, and same results hold for the commitment bad agent.

In EV, as the agent lives for an infinite number of periods, $m = (n - i)$ is not feasible. Hence, the good agent holds strict reputation incentives in any period where the bad agent holds reputational incentives. As $n \rightarrow \infty$, the bad agent has reputational incentives in each period. The agent is never hired in the equilibrium.

In the finitely repeated interaction as defined in this chapter, there is some period i after which the agent does not have reputational incentives. The agent is hired after such a period as long as the reputation does not fall below $1 - 2k$. $i = \max \{1, \lceil n - t + 1 - 1/k \rceil\}$ is the period where the hiring starts. There are no strict reputational incentives left, and both types of agents play stage-game maximizing actions in any period where hired. I have the same equilibrium with the commitment bad type. The formal definition of the equilibrium is provided below.

Proposition 1.4. *Let $P(t) \leq \rho_1 < P(t + 1)$, the following assessment constitutes the agent-optimal perfect Bayesian equilibrium. The first period the agent is hired is j , such that $j = \max \left\{ 1, \left\lceil n - t + 1 - \frac{1}{k} \right\rceil \right\}$ ¹. The agent will be hired until the end of the game after any history in which action 0 is observed. If 0 is never observed, the agent is hired for a total of t periods. $\mu_i = 0$, & $\nu_i = 0$ for all the periods the agent is hired.*

Proof. Appendix A. □

Lower initial reputation implies higher reputational incentives. Therefore, the agent has a stronger incentive to play action 0 in state 1. Thus, she will not be hired in the initial periods. For every period the agent is not hired, there are fewer periods left to be played. The lower the number of remaining periods, the lower the reputational incentives are. Hence, the agent's reputational concerns decrease.

As the agent's reputation is not updated in the periods she is not hired, she will not have strict reputational incentives left after a certain number of periods without hiring. The short-term principal's objective is to maximize one period expected

¹ $\lceil x \rceil$ is equal to the lowest integer greater than or equal to x

payoff; hence the principal starts to hire the agent as soon as she has no strict reputational incentive.

$k < 1/2$ implies $n - \max \left\{ 1, \left\lceil n - t + 1 - \frac{1}{k} \right\rceil \right\} \geq t$. Therefore, in period j , where the agent is first hired, there will be at least t periods left until the project's end. If action 0 is observed in any period, then the good agent's type will be fully revealed, and she will be hired for the remainder of the game. On the other hand, if no action 0 is observed in any period, then the agent is hired for exactly t periods. Bayes' rule implies that starting from period j , if action 1 is observed for all t periods, the reputation falls below $1 - 2k$, and the future principals do not hire the agent.

Suppose $P(6) \leq \rho_1 < P(7)$ and fix $n = 10$ and $k = 1/3$, then:

$$\max \{1, \lceil n - t + 1 - 1/k \rceil\} = 2$$

Suppose for the sake of contradiction that the agent is hired in the first period. Consider the following case where $\theta_i = 1 \forall i \in 1, 2, \dots, 6$. If $\mu_i = 0 \forall i \in 1, 2, \dots, 6$, then Bayes' rule implies that $\rho_7(1 | a_i = 1 \forall i < 7) < 1 - 2k$. Hence, the agent will not be hired in period 7. Now consider period 6. If agent plays action 0, then she gets a payoff of k . On the other hand, as $\rho_7(0) = 1$, the agent gets the payoff of $k - 1 + 4k$ if she deviates to playing action 0. Hence, $k = 1/3$ implies that $\mu_6 = 1$ after such history where action 1 is observed in all the previous periods. Next, consider the period 6 principal. $\mu_6 = 1$ implies that the principal gets a negative payoff if he hires the agent. Hence, the principal does not hire the agent in period 6. Continuing with backward induction, $\mu_5 = 1$ if the agent is hired in period 5. The reason is that the agent will not be hired if $a_5 = 1$ is observed. Thus, period 5 principal does not hire the agent either. Continuing in this fashion, I see that the agent is not hired in the first period. More precisely, the hiring of agent does not start as long the agent has strong reputation concerns. Hence, hiring starts in the second period in this case, where there are no more reputation concerns left. Bayes' rule implies that if $\mu_i = 0 \forall i$, then the agent is hired for 6 periods if action 1 is observed in each period. Hence, the agent gets an expected payoff strictly greater than $6k$. On the

other hand, $\mu_i > 0$ in any period lead to a payoff weakly less than $9k - 1$ (as hiring starts in the second period). $k = 1/3$ implies that the agent has no incentive left to play costly action 0 if hiring starts in the second period. The agent is hired in periods $\{2, 3, 4, 5, 6, 7\}$ irrespective of the action observed. On the equilibrium path, the lowest reputation level in period 7 is in the range $[1 - 2k, (1-2k)/(1-k))$ which is attained when action 1 is observed in all the periods. The agent is hired in $i = 7$ as $\rho_7 \geq 1 - 2k$. However, if action 1 is observed in period 7, then Bayes rule implies $\rho_8(1) < 1 - 2k$. Thus, the agent is not hired for the remaining three periods.

Remark 1.1. If $P(t) \leq \rho_1 < P(t+1)$ and $n - t \leq \frac{1}{k}$, then

$$\max \left\{ 1, \left\lfloor n - t + 1 - \frac{1}{k} \right\rfloor \right\} = 1$$

Note that this condition implies the agent has no strict reputational incentive if:

1. Initial reputation is sufficiently high. Higher initial reputation implies higher t . (and/or)
2. k is sufficiently small. For k close enough to 0, there is no reputational concern for almost all reputation levels. This is a natural outcome as lower k means a lower stage game payoff, hence lower reputational incentives. (and/or)
3. n is sufficiently small. If the agent's lifespan is short, she will not have reputation concerns.

Remark 1.2. If the initial reputation is such that $P(t) \leq \rho_1 < P(t+1)$ and $t = n-2$ or $n-1$, $\max \left\{ 1, \left\lfloor n - t + 1 - \frac{1}{k} \right\rfloor \right\} = 1$. Hence the agent is hired starting from the first period. If $P(n-1) \leq \rho_1 < P(n)$, then the agent is hired starting from the first period and will be hired up until the last period if action 0 is never observed. In the history h_n where no action 0 is observed, then $\rho_1 < P(n)$ implies $\rho_n < 1 - 2k$,

hence $\lambda_n = 0$. If action 0 is ever observed, the agent is hired for the remainder of the game. Similarly, if $P(n-2) \leq \rho_1 < P(n-1)$, then the agent is hired for the first $n-2$ periods in the history where no action 0 is observed. In such history $\rho_{n-1} < 1-2k$ by Bayes' rule.

There is a range of the initial reputation such that the agent has no reputational incentive and is hired for any k . As the agent's life span is finite, there is always a period in which sufficiently fewer periods are left to be played. This phenomenon contrasts with EV's setup, where the agent is never hired irrespective of the initial reputation. The crucial difference between our model and EV is that our model is finite. In fact, when $n \rightarrow \infty$, then $P(t) \rightarrow 1$. Thus, there will be a strict reputational incentive for all initial reputation levels less than 1, which will lead the principal never to hire the agent. Hence, our results hold with EV when $n \rightarrow \infty$.

1.5 Long-term principal

When the principal is also a long-term player, the agent may also be hired in the periods with reputational incentives. The reasoning is that the long-term principal benefits from the increase in reputation in the future periods. In this section, I will define equilibrium where both types of agents are strategic. Later, I will provide the analysis for the commitment bad agent.

As in the previous section's analysis, if the number of periods is below a certain level, then the agent has no reputational concern in the unique equilibrium. More precisely, if $n < \frac{1}{k} + 1$, then $\mu_i = 0 = \nu_i$ for all periods. Hence, the results apply to the strategic and commitment bad agent analysis. As $k < 1/2$, $n < \frac{1}{k} + 1$ is directly satisfied for $n \leq 3$. If the interaction lasts for three or fewer periods, then there are no reputational incentives.

Let $t \in \{1, 2, \dots, n-1\}$ and define

$$R(t, n) = \frac{(1-2k)2^{t-1}}{(1-2k)2^{t-1} + (n-t+2)k}$$

If the initial reputation is sufficiently high, i.e., $\rho_1 \geq R(n, n)$, then the agent is hired every period irrespective of the history in the game. Consider period n . The lowest reputation level of the agent is attained in the history in which action 1 is observed in all the previous periods. Bayes' rule implies that if $\rho_1 \geq R(n, n)$ then $\rho_n \geq 1 - 2k$ even after such history. Hence, the principal hires the agent in any period. This implies that the agent has no incentive to take a costly reputation action, i.e. $\mu_i = \nu_i = 0 \forall i \in N$.

Next, consider middle initial reputation levels. The following lemma will be used in the further results.

Lemma 1.6. *Take any perfect Bayesian equilibrium and suppose that $\mu_i = \nu_i = 0$ for $i = j, j+1, \dots, n$. Then the agent is hired in period j if and only if $\rho_j \geq \frac{1 - 2k}{1 + (n - j)k}$.*

Proof. Appendix A. □

When the principal makes the hiring decision in period j , he hires the agent if the agent's reputation is above a certain level so that the principal expects a non-negative payoff. The principal's expected payoff depends on whether the agent will have reputation concerns in the remaining periods. Lemma 1.6 states the lowest reputation level for which the principal hires the agent if there will be no reputation

concerns in the future. It corresponds to $\rho_j \geq \frac{1 - 2k}{1 + (n - j)k}$.

Suppose that $\mu_i = \nu_i = 0$ for $i = 1, 2, \dots, t$. Bayes' rule implies, if the initial reputation is such that $R(t, n) \leq \rho_1 < R(t + 1, n)$ for some $t \in \{1, 2, \dots, n - 1\}$,

then in the history where no action 0 is played, $\rho_t \geq \frac{1 - 2k}{1 + (n - t)k}$ and $\rho_{t+1} <$

$\frac{1 - 2k}{1 + (n - t - 1)k}$. By Lemma 1.6, it implies the principal hires the agent in period t if he expects no reputation incentive. In addition, if action 1 is observed in period t , then the principal does not hire the agent in period $t + 1$.

In the equilibrium path where there are no reputation concerns in any period, $R(t, n) \leq \rho_1 < R(t + 1, n)$ implies that the agent is hired for exactly t periods if no action 0 is observed. If action 0 is observed in any period, then the agent is hired

for the rest of the relationship.

So far, I have defined the equilibrium behaviors with no reputation concerns. Next, I will show the necessary and sufficient conditions for the absence of reputation concerns in any equilibrium. Define *critical period* i as a period in which $\lambda_{i+1}(0) = 1$ and $\lambda_{i+1}(1) = 0$. If $a_i = 1$ is observed, then the agent is not hired for the remaining periods. On the other hand, if $a_i = 0$ is observed, she is hired for 1 or more periods in the future. Consider critical period t . The sequential rationality constraint of the agent implies that if $k > (n - t)k + k - 1$, then the agent has no incentive to play the costly reputation action. Now consider period $t - 1$. Good and bad agents' sequential rationality constraints are as follows:

$$\text{good} : 2k + \frac{1}{2}(n - t)k > k - 1 + k + (n - t)k$$

$$\text{bad} : 2k > k - 1 + k + (n - t)k$$

Realize that the absence of reputation concerns in the critical period implies that there are no reputation concerns in the previous periods. Hence, if $R(t, n) \leq \rho_1 < R(t + 1, n)$ and $n < \frac{1}{k} + t$, then $\mu_i = 0$ and $\nu_i = 0 \forall i \in N$ in any equilibrium. The principal hires the agent for the first t periods if no action 0 is observed. Proposition 1.5 formalizes the argument.

Proposition 1.5. 1. If $\rho_1 \geq R(n, n)$, then the agent has no reputational incentive in the agent-optimal perfect Bayesian equilibrium, i.e., $\mu_i = \nu_i = 0 \forall i \in N$, and she is hired in every period irrespective of the history.

2. If there exists $t \in \{1, 2, \dots, n - 1\}$ such that $R(t, n) \leq \rho_1 < R(t + 1, n)$ and $n < \frac{1}{k} + t$, then the agent has no reputational incentive in the agent-optimal perfect Bayesian equilibrium, i.e., $\mu_i = \nu_i = 0 \forall i \in N$, and she is hired in periods $1, 2, \dots, t$ after any history. In period $i \geq t + 1$, she is hired if and only if action 0 is observed in some period $j < i$.

3. If $\rho_1 < R(1, n)$ and $n < \frac{1}{k} + 1$, then the agent has no reputational incentive in

the agent-optimal perfect Bayesian equilibrium, i.e., $\mu_i = \nu_i = 0 \forall i \in N$, and she is never hired.

Proof. Appendix A. □

In Proposition 1.1, I stated that if $n < \frac{1}{k} + 1$, then there are no reputation concerns in any equilibrium. Proposition 1.5 states that, $n < \frac{1}{k} + 1$ is not necessary but sufficient condition for the absence of reputation concerns in any equilibria. If the initial reputation of the agent is so that she will be hired for at least t periods, i.e. $R(t, n) \leq \rho_1 < R(t + 1, n)$, then the condition becomes $n < \frac{1}{k} + t$. Given initial reputation level of the agent, if n is small or k small, then the agent has no reputation concern. Note that, the absence of reputation concerns implies that the bad agent always plays action 1.

Reputational incentives are determined by n and k : Larger k and n lead to stronger reputational incentives. If n or k is small, there is no reputation concern. In other words, both the good and bad agents play their ideal actions. Therefore, there is no bad reputation in this case. Since the bad agent always plays action 1, if action 0 is observed in any period, the principal will hire the agent in every consecutive period. Whether the principal hires the agent in the first period depends on the initial reputation of the agent. If that reputation is large enough, i.e., $\rho_1 \geq R(1, n)$, the agent is hired in the first period. Hiring behavior after action 1 depends on the initial reputation as well. If $R(t, n) \leq \rho_1 < R(t + 1, n)$, for some $t \in \{1, 2, \dots, n - 1\}$, then the agent is hired for the initial t periods irrespective of the history.

The principal's hiring decision depends on the principal's expected continuation payoff, unlike the short-term principal. If the initial reputation is not less than $\frac{1 - 2k}{1 + (n - 1)k}$, then the principal has a non-negative expected continuation payoff. Hence, the agent is hired in the first period. If the initial reputation is below that level, then the agent is never hired. Note that this threshold is strictly less than $1 - 2k$, which is the lowest reputation level so that the short-term principal hires the agent.

Recall that, when the principal is a short-term player, the agent is hired t number of periods if no action 0 is observed if $P(t) \leq \rho_1 < P(t+1)$ where $P(t) = \frac{(1-2k)2^{t-2}}{(1-2k)2^{t-2} + k}$. Observe that $P(t) > R(t)$ and $P(t+1) > R(t+1)$. It implies that, for a given initial reputation level, the agent is hired for an equal or more number of periods where the principal is also a long-term player. Hence, when the principal is also a long-term player, the lower threshold for hiring the agent and the average number of hiring periods increase.

Next, I will analyze how equilibrium behavior and payoffs vary with the parameters of the game as long as $n < \frac{1}{k} + t$ holds.

1.5.1 What happens when n increases?

If the length of the interaction increases, the threshold for hiring, i.e., $R(1, n)$ decreases. Mathematically, it is easily seen as $R(1, n)$ is decreasing in n . The larger the interaction lasts, the smaller stake each period receives. Consider the worst scenario for illustration. The initial reputation is sufficiently small so that the agent is hired in the first period but is only hired in the second period if $a_1 = 0$ is observed. As the probability of the agent being bad is so high, the principal's expected first-period payoff is negative. However, in the absence of reputation concerns, $\rho_2(0) = 1$. Hence, with a probability of $\frac{1}{2}\rho_1$, the principal gets an expected payoff of k in the remaining periods. In the equilibrium path, only uncertainty about the agent's type lies in period 1. Thus, the principal hires the agent for the lower reputation levels as the stake of the first period, i.e., $\frac{1}{n}$, decreases.

Next, consider t . As n increases, the number of periods during which the agent is hired (even if the history has no action 0) also increases. The intuition is as follows. Take an equilibrium in which there are no reputational concerns and consider period j such that $\rho_j \geq \frac{1-2k}{1+(n-j)k}$ and $\rho_{j+1}(1) < \frac{1-2k}{1+(n-j-1)k}$. In words, period j is the critical period given that no action 0 is observed in the previous histories. If action 0 is observed in period j , then the agent is hired for the remaining periods.

If action 1 is observed, then the agent is never hired afterward. Higher n implies higher $n - j$; hence, there are more periods left to be played. Thus, increasing n implies the principal will hire for lower reputation levels in the critical period. This, in turn, implies that, for a given initial reputation level, the agent will be hired for an equal or more number of periods. Hence, $t(n)$ is weakly increasing in the number of periods.

The effect of change in the number of periods on the expected payoffs is ambiguous. The following has a positive effect on the continuation payoffs. As the number of interaction periods increases, t (the agent is hired for more periods) may increase. Moreover, $n - t$ may also increase or do not change. In any case, the average number of periods that the agent is hired increases. Also, if t increases, the probability that action 0 is observed also increases. The negative effect of the increase in n is that as the weights of the periods are normalized to add up to 1, each period's weight $1/n$ decreases. Despite the fact that the agent will be hired for more periods on average, the weight of each period will decrease. Hence, I can not conclude on the effect of period size on payoffs.

1.5.2 What happens when k increases?

The ex ante equilibrium payoffs of the players are as follows:

$$V_1^P = \frac{1}{n} \left(\rho_1 \left(nk - \frac{n}{2^t}k + \frac{1}{2}t - \frac{1}{2}tk \right) + tk - \frac{1}{2}t \right)$$

$$V_1^g = \left(1 - \frac{1}{2^t} + \frac{t}{2^t n} \right) k$$

The expected payoffs of both the principal and the good agent increase. It is intuitive as k denotes the value of the project. The threshold for hiring decreases. The same reasoning as subsection 1.5.1 applies. In the critical period, higher k implies a higher expected continuation payoff to the principal. As future expected payoff of the principal is higher in the critical period, he hires the agent for lower reputation levels, i.e. $\frac{1 - 2k}{1 + (n - i)k}$ is decreasing in k . Thus, fixing an initial reputation level,

increasing k parameter implies the agent will be hired for equal or more periods. Hence, t is weakly increasing in k .

If I treat k parameter as the project's size, I reach a natural conclusion as follows. The more significant (the more important) the project is, the higher the hiring frequency.

1.5.3 When there will be reputation concerns?

If n or k is large, then there are reputation concerns. Let us start with the complementary proposition to Proposition 1.5.

Proposition 1.6. 1. If there exists $t \in \{1, 2, \dots, n-1\}$ such that $R(t, n) \leq \rho_1 <$

$R(t+1, n)$ and $n \geq \frac{1}{k} + t$, then the agent has reputational incentives in any perfect Bayesian equilibrium, i.e., $\mu_i = \nu_i = 0$ can not hold $\forall i \in N$, .

2. If $\rho_1 < R(1, n)$ and $n \geq \frac{1}{k} + 1$, then the agent has reputational incentives in any perfect Bayesian equilibrium, i.e., $\mu_i = \nu_i = 0$ can not hold $\forall i \in N$, .

Proof. Appendix A. □

The intuition behind Proposition 1.6 is straightforward. Following the analysis in Proposition 1.5, given $R(t, n) \leq \rho_1 < R(t+1, n)$ and $n \geq \frac{1}{k} + t$, suppose for contradiction that the agent has no reputation incentive in the initial t periods. She will not be hired from period $t+1$ onward if no action 0 is observed. Having no reputation concern means the bad agent plays action 1 in any period. Hence, the bad agent gets a payoff of $\frac{t}{n}k$. If the bad agent deviates to action 0 in any period, she gets a payoff of $\frac{1}{n}(k-1+(n-1)k)$. $n > \frac{1}{k} + t$ implies that it is a profitable deviation. Hence, there will be reputation concerns for some period given $n > \frac{1}{k} + t$.

The extent of the reputation concerns depends on the relation between n , k and t . For illustration, proposition 1.7 demonstrates the equilibrium behavior for a specific case.

Proposition 1.7. *Let $\rho_1 < R(2, n)$ and $\frac{1}{k} + 2 \geq n > \frac{1}{k} + 1$, the following constitutes the principal and agent optimal equilibrium:*

$$\begin{aligned} \mu_1 &= 1, \nu_1 = \frac{1}{2^{n-3}} \frac{\rho_1 k}{(1 - \rho_1)(1 - 2k)}, \mu_i = 0 = \nu_i \forall i \in N/1 \\ \lambda_1 &= \begin{cases} 1, & \rho_1 \geq \frac{(1 - 2k)2^{n-3}}{(n-1)k(2^{n-2} - 1)} \\ 0, & \text{otw} \end{cases} \\ \lambda_2(0) &= 1, \lambda_2(1) = 0. \lambda_i(0) = 1 = \lambda_i(1) \forall i \in N/\{2, n\} \text{ and } \lambda_n(0) = 1. \\ \lambda_n(1) &= \begin{cases} \frac{1}{k} - (n-2), & a_i = 1 \forall i \in \{2, 3, \dots, n-1\} \text{ in } h_n, \\ 1, & \text{otw} \end{cases} \end{aligned}$$

Proof. Appendix A. □

The case analyzed in the Proposition 1.7 is a special case where the initial reputation is small and n is so that there are reputation concerns only in the first period. If the initial reputation is small, $\rho_1 < R(2, n)$, and both types of agents play the stage-game maximizing strategies, then the principal only hires the agent in the second period if action 0 is observed. $n > \frac{1}{k} + 1$ implies the bad agent will deviate to play action 0. Hence, no reputation concern in the first period can not hold in any equilibrium. Moreover, $\frac{1}{k} + 2 \geq n$ implies there are no more reputation concerns from period 2 onward. I have a discrete number of periods. I need to introduce a mixing strategy for the principal so that the agent can be indifferent between both actions. Otherwise, reputation concerns imply $\mu_i = 1 - \nu_i$. There will be no update on the reputation level. However, $\lambda_2 = 0$ if $\rho_2 = \rho_1$ as the initial reputation is sufficiently small. Hence, strategies should be so that there is an update in the reputation level, which necessitates indifference for at least one type of agent.

Note that the threshold initial reputation for hiring, $\frac{(1 - 2k)2^{n-3}}{(n-1)k(2^{n-2} - 1)}$ is strictly less than $\frac{1}{2} - k$. It implies there is a range of initial reputation such that the agent

is also hired for some $\frac{(1-2k)2^{n-3}}{(n-1)k(2^{n-2}-1)} \leq \rho_1 < \frac{1}{2} - k$. In words, if the agent has reputation concerns, then she is hired for lower initial reputation levels compared to the case with the absence of reputation concerns. As the equilibrium behaviors vary with n , k and ρ_1 , I can not provide a general pattern.

1.5.4 Commitment bad agent

In this subsection, I will define the equilibrium for the interaction lasting sufficiently long with the commitment bad agent's presence. More precisely, reputation incentives emerge in the interactions lasting more than $(1/k + 1)$ periods. From Proposition 1.1, I know the equilibrium behavior in the complementary case. Parallel to the short-term principal's analysis, the agent has a reputational incentive in the initial periods. After a certain period, more precisely, the period $i > n - \frac{1}{k}$, the agent has no reputational incentives for the remaining periods. When the principal is a short-term player, he avoids hiring the agent in periods with reputational incentives. However, the long-term principal benefits from the increase in the agent's reputation in the future. Thus, the agent is also hired in the initial periods.

Another difference with the short-term principal is that the agent is hired for lower reputation levels. If the principal is short-term, his objective is to maximize the expected stage game payoff. Thus, the agent is not hired if the expected stage-game payoff is negative. On the other hand, the long-term principal tries to maximize the expected continuation payoff. He may have a positive expected continuation payoff even if he receives a negative payoff for some periods. In such instances, the agent is hired. In that sense, "bad reputation" incentives do not hurt the principal much. The losses stemming from the principal being a short-term player can be avoided.

Further, I show that the (good) agent prefers to postpone screening as much as possible. The reasoning is straightforward. Consider a period in which the agent has reputational incentives. If action 0 is observed, she is hired for the remainder of the game. If action 1 is observed, she is hired for some periods, say l periods, as long as no action 0 is observed. The good agent's sequential rationality constraint

implies that $\mu_i = 1$ may be an equilibrium strategy if l is small enough. She screens her type and is hired forever. However, she can increase her expected continuation payoff by postponing the full screening by one period. If she does not screen her type in the current period, there is a $1/2$ probability that $\theta_{i+1} = 0$, and she will reveal her type costlessly in the next period. Thus, as long as the agent's reputation is sufficiently high so that she will be hired after action 1, she does not play action 0 with a positive probability (given the state is 1) in the equilibrium.

Consider any period $i < n - \frac{1}{k}$. $n \geq \frac{1}{k} + 1$ implies that there will be reputation concerns. In any such period, the principal's expected payoff is non-negative if $\rho_i \geq \frac{1-2k}{2(n-i)k}$. Thus, the agent is only hired for $\rho_i \geq \frac{1-2k}{2(n-i)k}$.

1. If $i+1 \leq n - \frac{1}{k}$, then the agent will have reputation concerns in the next period.

Hence, the principal will hire the agent in period $i+1$ if $\rho_{i+1} \geq \frac{1-2k}{2(n-i-1)k}$.

- (a) If ρ_i is such that $\rho_{i+1}(1) < \frac{1-2k}{2(n-i-1)k}$, i.e. the principal does not hire the agent after action 1, then the agent plays action 0 with probability of 1 and screens her type.

- (b) If ρ_i is such that $\rho_{i+1}(1) \geq \frac{1-2k}{2(n-i-1)k}$, i.e. the principal hires the agent after action 1, then the agent plays stage-game maximizing strategy. As noted before, $\mu_i = 0$ as long as $\lambda_{i+1}(1) = 1$.

2. If $i+1 > n - \frac{1}{k}$, then the agent will no more have reputation concerns in the next periods. The principal's expected payoff in period $i+1$ is non-negative if $\rho_i \geq \frac{1-2k}{1+(n-i)k}$. Thus, the agent is only hired for $\rho_i \geq \frac{1-2k}{1+(n-i)k}$.

- (a) If ρ_i is such that $\rho_{i+1}(1) < \frac{1-2k}{1+(n-i-1)k}$, i.e. the principal does not hire the agent after action 1, then the agent plays action 0 with probability

of 1 and screens her type.

- (b) If ρ_i is such that $\rho_{i+1}(1) \geq \frac{1-2k}{1+(n-i-1)k}$, i.e. the principal hires the agent after action 1, then the agent plays stage-game maximizing strategy. As noted before, $\mu_i = 0$ as long as $\lambda_{i+1}(1) = 1$.

In any period $i > n - \frac{1}{k}$, there are no reputation concerns, i.e. $\mu_i = 0$. Proposition 1.8 formalizes the equilibrium.

Proposition 1.8. *If $n \geq \frac{1}{k} + 1$, the unique perfect Bayesian equilibrium strategies in period i is as follows:*

1. For the initial periods such that $i < n - \frac{1}{k}$:

$$\lambda_i = \begin{cases} 0, & \rho_i < \frac{1-2k}{2(n-i)k} \\ 1, & \rho_i \geq \frac{1-2k}{2(n-i)k} \end{cases}$$

$$\mu_i = \begin{cases} 0, & \rho_i \geq \rho'_i \\ 1, & \rho_i < \rho'_i \end{cases} \text{ such that } \rho'_i = \begin{cases} \frac{2-4k}{2(n-i-2)k+1}, & i \leq n-1-\frac{1}{k} \\ \frac{2-4k}{2+(n-i-3)k}, & i > n-1-\frac{1}{k} \end{cases}$$

2. For $i > n - \frac{1}{k}$:

$$\lambda_i = \begin{cases} 0, & \rho_i < \frac{1-2k}{1+(n-i)k} \\ 1, & \rho_i \geq \frac{1-2k}{1+(n-i)k} \end{cases}$$

$\mu_i = 0$.

Proof. Appendix A. □

The principal's hiring decision in any period depends on the expected continuation payoff. It, in turn, depends on whether the agent will have a reputational

concern in the period. More precisely, there are reputational incentives in periods $i \leq n - \frac{1}{k}$. In the periods where the agent has reputational incentives, the agent chooses to invest in her reputation if she believes that she will not be hired after action 1. If the reputation is below ρ'_i , then the agent knows that $\rho_{i+1}(1) = \frac{\rho_i}{2 - \rho_i}$ is not sufficiently high for the principal to hire the agent in period $i + 1$ given $a_i = 1$. Thus, the agent reveals her type, i.e., $\mu_i = 1$.

If the initial reputation is sufficiently high, then the agent never has a reputational incentive. More precisely, Bayes' rule implies that, if $\rho_1 \geq \frac{2^{n-1/k}(1-2k)}{2^{n-1/k}(1-2k)+1}$, then $\rho_{n-1/k} \geq \frac{2-4k}{3-4k}$. In other words, the agent is hired in the first $n - 1/k + 1$ periods irrespective of the history. Hence, $\mu_i = 0$ for all $i \in N$. The principal hires the agent in any period i as long as $\rho_i \geq \frac{1-2k}{1+(n-i)k}$. On the other hand, $\rho_1 < \frac{2^{n-1/k}(1-2k)}{2^{n-1/k}(1-2k)+1}$ implies by Bayes' rule that if no action 0 is observed in the first i periods for some $i \leq n - 1/k$, then $\lambda_{i+1} = 0$. In such period, the agent has an incentive to play action 0 with a positive probability.

1.6 Conclusion

This chapter provided an analysis of a finitely repeated principal-agent interaction. As expected, I showed that the agent has reputation concerns depending on the model's parameters. Firstly, reputation concerns increase with the length of the interaction. The longer the lifespan of the agent (or the interaction, in the case of long-term principal), the higher the potential payoff the agent can earn in the future. Hence, the long-term incentives to build a high reputation may triumph over the short-term incentives to play the stage-game optimal action. On the extreme, if the number of periods tends to infinity, both types of agents have reputation concerns in any period. Thus, the principal never hires the agent, as shown by Ely and Välimäki [2003]. Nevertheless, if the agent (or the interaction) has a finite lifespan, there is always a period after which there are no reputation concerns.

Second, reputation concerns decrease in k , the value of the project. Note that a change in k has two effects. Decreasing k implies: (a) increasing (the absolute value of) the cost of reputation building ($k - 1$), and (b) decreasing the future expected payoff. Hence, reputation-building behavior is closely related to the value of the project. Thirdly, the strength of the reputation incentives is negatively correlated with the reputation level of the agent. The agent with a lower reputation level has higher reputation incentives, *ceteris paribus*.

The short-term and long-term principals act differently as expected. Short-term principal only focuses on the stage game payoff; hence, he avoids agents with reputation concerns. On the other hand, long-term principal's behavior is significantly different. Firstly, the long-term principal hires the agent even in periods with strict reputation concerns. The reason is that the long-term principal may benefit from the reputation concerns of the agent. What is called "bad reputation" by Ely and Välimäki [2003] becomes a valuable tool for the principal in our setup. It is due to "sorting effect" of reputation concerns [Morris, 2001]. The principal is able to observe the agent's actions and update his belief. Secondly, the long-term principal hires the agent for both lower initial reputation levels and weakly more periods on average. Hence, I conclude that the harm of reputation incentives can be avoided if both players are long-term.

Chapter 2

ENDOGENOUS STAKES IN A REPEATED PRINCIPAL-AGENT INTERACTION

2.1 Introduction

In this chapter, I introduce endogenous stakes to the analysis of Chapter 1. I ask the following questions: Suppose the agent, the principal, or a social planner can choose the size of the stakes involved in their relationship at every stage of the game. Would they choose them in a way to prevent a bad reputation, i.e., let the good and the bad agents separate at the beginning of the game? What is the equilibrium path of the size of the stakes? Does the principal, the social planner, or the agent choose to start the relationship large or small?

Since the bad reputation may hurt the good agent and the principal, one would think they would choose the stakes to eliminate reputational incentives. This can be achieved by starting the relationship large and assigning smaller stakes to the future. This intuition is correct in specific environments. In particular, the agent prefers to start large. When the agent chooses the stakes, she puts high stakes to initial periods. Hence, small stakes are left to the future periods, and reputation concerns are prevented.

However, I show that in some instances, intuition does not work. One critical case is as follows. Suppose a long-term principal faces an agent that is a strategic bad type with sufficiently high probability. In that case, he chooses to start the relationship small and gradually increases the size of the stakes. Furthermore, the equilibrium behavior induced by this choice of stakes involves a bad reputation. The good agent always plays the action contributing to her reputation while the bad agent completely mixes until the last two periods. This result is surprising even

more because this is also the equilibrium most preferred by the principal, i.e., he chooses to use a bad reputation to his advantage if the bad agent is strategic and the initial reputation is bad enough. Moreover, I show that the principal does not start large in any equilibrium.

Several empirically relevant scenarios, some of which have already been discussed by EV, fit one of the above environments. Lawyers might need a large case to kick start their careers. If a lawyer starts with minor cases, the client may believe that the lawyer will be more interested in increasing her reputation than the case. By starting with large cases, the lawyer can signal to the potential client that she will not have reputational incentives. On the other hand, a government official appointing judges (or attorney general assigning cases to state attorneys) may prefer to start with small cases. More minor cases in the initial periods minimize the harm of the reputational incentives. My analysis allows us to make equilibrium predictions in any such environment. I need to note that our results should be interpreted carefully since I am not considering other possible determinants of the allocation of stakes over time, such as learning about the agent's ability, learning on the job, etc.

I start with the short-term principal in Section 2.4, and then move to the long-term principal in Section 2.5. I show that if the agent sets stakes, she prefers to start the interaction large and decrease the stakes gradually. There are instances when the agent holds no reputation concern in the equal-stakes model. For example, the interaction length may be short, or the value obtained from the project may be small. In that case, the agent's ability to design the career path benefits her in two aspects. First, she obtains a higher payoff by allocating the high stakes to the periods where she is hired for sure. Second, she is also hired for lower reputation levels (in which she is not hired in the equal-stakes model). In the complementary case, the agent has reputation concerns in the equal-stakes model (It occurs if the interaction lasts for many periods or the value of the project is large). By starting large, the agent leaves sufficiently small stakes for future interaction so that she has no incentive to invest in her reputation. Therefore, the "bad reputation" outcome can be avoided by changing the allocation of relative stakes. The agent obtains a

higher payoff if she is hired.

However, the agent's ability to set the stakes does not always benefit her. Consider the following case. The agent has a small initial reputation level and holds reputation concerns in the equal-stakes model. The principal expects a non-negative payoff, hence, hires the agent. However, the principal may not hire the agent if the agent sets the relative stakes. The agent can not commit to a particular allocation of stakes for all periods at the beginning of the interaction. Once the agent's reputation level increases above a certain level, the agent will put the highest possible stakes in every successive period. Mostly, it corresponds to the allocation that yields a zero continuation payoff to the principal. As the principal gets a negative payoff in the earlier periods with reputation incentives, it means a negative continuation payoff for the principal in the first period. Hence, the principal never hires the agent.

When the long-term principal sets relative stakes, he never prefers to start large. I show that if the bad agent is a commitment type, then the principal starts the interaction small so that the good agent reveals her type in the first period. On the other hand, if the bad agent is also a strategic player, the principal never sets the stakes such that the agent would strictly prefer to play the stage-game maximizing action. The principal sets the stakes so that at least one type of agent is indifferent between playing both actions. Endogenous stakes improve the principal's payoff, and the agent is hired for lower reputation levels. Additionally, I show that if the agent's initial reputation is sufficiently bad, the principal starts small and increases the stakes gradually. The bad agent is indifferent between the two actions in each period. The good agent's type is not revealed until the last period in the equilibrium. The principal sets stakes for each period so that he can manipulate the reputation incentives of the agent and get the maximum expected value from the project. It is consistent with the reputation literature on gradualism. Watson [1999, 2002a], Andreoni and Samuelson [2006] are the most prominent ones. Compared to the setup where each period receives equal weight, the principal's ability to set the relative stakes enables him to hire the agent with lower initial reputation levels.

Finally, I introduce a social planner who wants to maximize social welfare, which

is defined as the summation of the short-term principal's payoffs across all periods. The social planner can be thought of as an owner of a legal advice bureau. The owner has lawyer employees who work for short-term principals. I show that if the bad agent is a commitment type, the social planner starts the agent's career path with small stakes and gradually increases the relative stakes. The social planner sets stakes so that the good agent is indifferent in each period.

On the other hand, if the bad agent is also a strategic player, the social planner allocates equal stakes to the initial periods, serving as the evaluation phase. The bad agent plays his favorite action in each period. The length of the evaluation phase depends on the agent's initial reputation: the higher the initial reputation is, the longer the evaluation phase is. If the good agent does not reveal her type by the end of the evaluation phase, she is not hired after that. On the other hand, if she reveals her type, the evaluation phase ends, and she is hired for the remainder of the interaction.

2.2 Literature

This chapter belongs to the literature on gradualism and dynamic allocation of stakes, which have been studied in various frameworks. It is widely accepted that starting the interaction with small stakes and increasing them over time leads to maximum payoffs. The pioneering papers on gradualism are those by Watson [1999, 2002a]. He studied infinitely repeated prisoners' dilemma game where the partners jointly and dynamically determine the periods' relative weights. There is two-sided incomplete information about the types of players. A low type prefers to betray, whereas a high type chooses to cooperate as long as the other player cooperates. Both of Watson's papers conclude that if the interaction starts small, cooperation is viable, and there is an equilibrium where high types cooperate perpetually.

Another paper on gradualism is by Andreoni and Samuelson [2006]. They experimented with a twice-played prisoner's dilemma game. They varied the relative stakes of the periods and found that the stakes' optimum allocation is approximately one-third to two-thirds. They experimentally verified that starting small is

the best concerning the total payoff of the players. Andreoni et al. [2019] provided an extension to that analysis by endogenizing the relative stakes' decision. In their experiments, players chose the relative stakes, and they found that starting small increases cooperation and social payoff. Grosskopf and Sarin [2010] provided another experimental work where a long-run player faced short-run players. They found that the reputation is seldom harmful, and its beneficiary effects are not tremendous as thought.

Atakan et al. [2020] studied a repeated interaction between a sender and a receiver for a finite number of periods. In their model, the sender is informed and interacts with an uninformed and possibly bad receiver. The sender holds information and decides whether or not to share accurate information with the receiver in each period. The sender or receiver sets the importance parameter in each period in this framework. With its finite nature and dynamically set relative stakes, Atakan et al. [2020] provide another starting point for our paper. They showed that it is payoff-dominant for both receiver and the sender to start the interaction small and increase gradually through time.

2.3 The Model

I have same model as in Chapter 1. Hence, in this section, I focus on how stakes are set endogenously. In each period i , a player sets the relative stake of the period i decision. More precisely, in period i , $\delta_i \in [0, 1]$ is the weight of all the future periods $i + 1, i + 2, \dots, n$ while $1 - \delta_i$ is the weight of period i . In other words, in each period i , the player sets the proportion of remaining decisions to be made in that period as $(1 - \delta_i)$ and the proportion to leave to subsequent periods as δ_i . Thus, if I define a weight for each period, γ_i , then $\gamma_i = (1 - \delta_i) \prod_{j=1}^{i-1} \delta_j$.

I assume $\delta_i \in [\underline{\delta}, \bar{\delta}] \ \forall i$ such that $0 < \underline{\delta} = 1 - \bar{\delta} < \frac{k}{1+k}$. The assumptions $\underline{\delta} > 0$ and $\bar{\delta} < 1$ guarantee that each period receives a positive stake. I assume $\underline{\delta} < \frac{k}{1+k}$

so that $\underline{\delta} < \frac{1}{1+k} < \bar{\delta}$. The reason is that if δ is less than $\frac{1}{1+k}$, then the agent has no reputational incentive in any equilibrium. Likewise, if $\delta \geq \frac{1}{1+k}$, then there are reputational incentives in some equilibrium. Thus, I set a lower threshold for δ less than $\frac{1}{1+k}$ so that there will be no reputation incentives if $\delta_i = \underline{\delta}$. On the other hand, for the upper threshold to be meaningful, I assume $\bar{\delta} > \frac{1}{1+k}$, i.e. there may be reputational incentives for $\delta_i = \bar{\delta}$. Finally, I make the assumption $\bar{\delta} = 1 - \underline{\delta}$ to simplify the equilibrium analysis.

For any $i \in N$, let H_i be the set of all histories before decision i is made, i.e. $H_i = (o_1, o_2, \dots, o_{i-1})$ where $o_i = (\delta_i, \lambda_i, a_i)$. Hence, in period i , the players observe the previous actions, previous stake choices and δ_i before making the hiring decision. $\alpha_i \in \{\emptyset, \{0, 1\}\}$ denotes the outcome of the period i action. If the agent is not hired in period i , then $\alpha_i = \emptyset$. The strategy of the period i principal is given by $\Lambda_i : H_i \rightarrow \{0, 1\}$ where $\Lambda_i(h)$ is the principal's choice of λ_i . Principal's belief about the type of agent is defined by the probability of agent being good in period i : $\rho_i : H_i \rightarrow [0, 1]$. The principal observes all the previous actions of the agent but does not observe the true state in these periods. The agent moves after histories of this type $I_i = (h, \lambda_i, \theta_i)$ where I_i is the period i information set of the agent. The good agent's (mixed) strategy is given by $\mu_i = \text{prob}(a_i = 0)$ and the bad agent's strategy is given by $\nu_i = \text{prob}(a_i = 0)$. $\mu_i \& \nu_i : \mathfrak{I}_i \rightarrow \Delta[0, 1]$ where $\mathfrak{I}_i$ is the set of all period i information sets.

I define equilibria for both commitment and strategic bad types. Commitment bad type agent always plays stage-game maximizing strategy, which corresponds to $a_i = 1 \forall i \in N$. Note that $\nu_i = 0 \forall i \in N$ if the bad agent is a commitment type. The strategic bad type is a payoff-maximizing player who does not necessarily play the optimal stage-game action.

The timing of the stage game in period i can be summarized as follows:

1. The agent, the social planner, or the principal set δ_i .

2. The principal makes the hiring decision. If $\lambda_i = 0$, then the game moves to the next period. Both parties get a payoff of 0.
3. If $\lambda_i = 1$, nature chooses the state $\theta_i \in \{0, 1\}$.
4. The agent observes θ_i and takes the action $a_i \in \{0, 1\}$.

The stage game payoffs and stage game strategies are same with Chapter 1, given by the Tables 1.1 and 1.2. The expected continuation payoff of the agent type $A \in \{g, b\}$ (good and bad agent, respectively) in period i is given as follows: $V_i^A = \sum_{j=i}^n \gamma_j U_j^A$ where U_i^A denotes the stage game payoff. Period i short-term principal interacts with the agent for one period, he aims to maximize U_i^P .

Results of Lemmas 1.1, 1.2 and 1.3 continues to hold. In words, while playing action 0 weakly increases the agent's reputation, playing action 1 decreases it in any perfect Bayesian equilibrium. Hence, the good agent has no incentive to play action 1 when the state of the world is 0.

I continue to focus on the perfect Bayesian equilibria with Markovian property. In that regard, previous history (including δ_i) choices matter only in terms of their effect on the agent's reputation. Thus, if the agent has the same reputation level after two different histories, both the agent's and the principal's actions are the same in both cases. The decisions depend on the past observations only through their effect on the agent's reputation.

Markovian Property. For any $i \in N$ and $h, h' \in H_i$: $\rho_i(h) = \rho_i(h')$ implies $\lambda_i(a_i|h) = \lambda_i(a_i|h')$, $\mu_i(h) = \mu_i(h')$ and $\nu_i(h) = \nu_i(h')$.

When there are multiple equilibria, I use the agent or the principal optimality as the equilibrium selection criteria. The agent (principal) optimal equilibrium is the equilibrium that leads to the highest expected continuation payoff for the agent (principal) among all equilibria.

2.4 Short-term Principal

For the short-term principal, lower δ_i means higher γ_i , hence the higher the stakes at hand. As the principal is an expected payoff maximizer, fixing a strategy for the agent, the choice of δ_i does not affect the principal's hiring decision. However, δ_i has an indirect effect on the hiring decision as follows. If δ_i is sufficiently high, it implies the future is sufficiently more important than period i . The agent may have reputational concerns depending on the current reputation level, leading the principal not to hire the agent. Knowing that the agent does not benefit from setting high δ_i because the principal will not hire the agent if the agent plays action 0 with a high probability. Moreover, the agent prefers the equilibrium in which she does not have to screen her type.

2.4.1 Agent sets stakes

Define $\mu'_i = \frac{2k + \rho_i - 1}{\rho_i}$ from Section 1.3.1, which is the highest probability for which the agent plays action 0 and is hired in period i . I define equilibrium when the bad agent is a strategic player. I show that no type has a reputational incentive in the agent-optimal equilibrium, and the results hold for the commitment bad type.

In Chapter 1, Proposition 1.1, I showed that the agent has no reputational concern in the last three periods (or if $n \leq 3$) in any equilibrium. However, when the relative stakes can vary, there exists an equilibrium for which the agent has a reputational concern in any period except the period n . Consider period $n - 1$. If each period has the same stakes, then the agent will never take the cost of investing in the reputation only for k in the next period. But if δ_{n-1} is high enough, it implies the cost of reputation - $k - 1$ - is weighted by $1 - \delta_{n-1}$ and next period's payoff is weighted by δ_{n-1} which is sufficiently high. Thus, there exists ρ_{n-1} for which the agent has a reputational incentive for sufficiently high δ_{n-1} .

There exists a sufficiently high reputation level for which the agent has no reputation concern in any period, irrespective of the stakes in an equilibrium. The agent is hired forever even if action 1 is observed in all the periods. I show that the

agent prefers not to play the costly reputational action for this reputation level. It is no surprise because if it is possible for a party to survive for all periods by just maximizing their stage-game payoff, then they will not prefer any equilibrium where they take the costly reputational action. Thus, this equilibrium is the agent-optimal equilibrium.

The agent still prefers the equilibrium in which she does not take the costly reputational action for the lower reputation levels. She attains this by starting the career path with large stakes. There are two advantages of starting large. (1) If the reputation is not sufficiently good, there is a positive probability that the agent will not be hired after some period. Hence, putting high stakes on the periods where the agent is hired maximizes the expected continuation payoff. (2) Starting large implies fewer stakes left to the future, meaning the agent has no reputational incentive. As the period i principal knows that the agent has no reputational incentive, he hires the agent as long as $\rho_i \geq 1 - 2k$.

In any perfect Bayesian equilibrium, the bad agent mimics the good agent's choice of δ_i . Considering only pure strategy in δ_i , the reason is that, if the good agent's strategy in period i involves $\delta_i = \delta_i^g$, then the Bayes rule implies that $\rho_i(\delta_i|h_i) = 0$ for any $\delta_i \neq \delta_i^g$. If $\delta_i^b \neq \delta_i^g$ in equilibrium, then the bad agent is not hired and receives zero payoff. She can deviate to δ_i^g and obtain a positive payoff.

Recall that I defined $P(t) = \frac{(1-2k)2^{t-2}}{(1-2k)2^{t-2} + k}$ in Chapter 1 Section 1.4 where n denotes the total number of periods and $t \in \{1, 2, 3, \dots, n\}$.

Proposition 2.1. *If the agent chooses the stakes, the following assessment constitutes the agent-optimal perfect Bayesian equilibrium given $\rho_1 \geq 1 - 2k = P(1)$.*

$$\mu_i = 0 = \nu_i \text{ for all } i \in N.$$

1. Let $\rho_1 \geq P(n)$:

$$\delta_i \in [\underline{\delta}, \bar{\delta}] \quad \lambda_i = 1 \quad \forall i \in N$$

2. Let $P(t) \leq \rho_1 < P(t+1)$:

$\delta_i = \underline{\delta}$ as long as no action 0 is observed. The agent is hired for the total of

t periods if no action 0 is observed. If action 0 is observed in any period i , $\delta_j \in [\underline{\delta}, \bar{\delta}] \ \lambda_j = 1 \ \forall j \in \{i+1, i+2, \dots, n\}$.

Proof. Appendix B. □

In period 1, if the reputation is sufficiently high, i.e. $\rho_1 \geq \frac{(1-2k)2^{n-2}}{(1-2k)2^{n-2} + k}$, then the agent is hired for all the remaining periods irrespective of history. Thus, the agent has no reputational incentive for any allocation of stakes and sets relative stakes arbitrarily. For lower reputation levels, the reputation falls below $1-2k$ in some period $j < n$ in the history h_j where no action 0 is observed. Hence, the future expected continuation payoff of the agent in period 1 is less than $\delta_1 k$. The best allocation of stakes is to leave small stakes to the future. The interaction starts large, i.e. $\delta_1 = \underline{\delta}$. The agent has no reputational incentive for given stakes as $\underline{\delta} < \frac{1}{1+k}$, i.e. $\mu_1 = 0 = \nu_1$. If action 0 is observed, then $\rho_2(0) = 1$. The agent is hired for the remainder of the game, regardless of history. Otherwise, if $\rho_2 < 1-2k$, the agent is never hired for the remainder of the game. $\rho_2 \geq 1-2k$, the agent again sets $\delta_2 = \underline{\delta}$ with the same reasoning as in period 1. Hence, $\delta_i = \underline{\delta}$ for all i as long as no action 1 is observed.

By definition, $\delta_i = \underline{\delta}$ implies $\gamma_i = (1-\underline{\delta})\underline{\delta}^{i-1}$ for all $i \in N/n$ and $\gamma_n = \underline{\delta}^{n-1}$. Assuming $\underline{\delta} = 0.1$ and $n = 100$, Figure 2.1 represents the evolution of the stakes in the equilibrium.

2.4.2 Social planner sets stakes

In this section, I introduce a social planner where the long-term agent meets short-term principals. The objective is to model the interactions where a third party defines the terms of the interaction. The social planner can be thought of as a state-official organizing a particular career market to maximize the expected social outcome. I define this by the total of the short-term principal's payoffs. The social planner's objective is in parallel with that of the long-term principal.

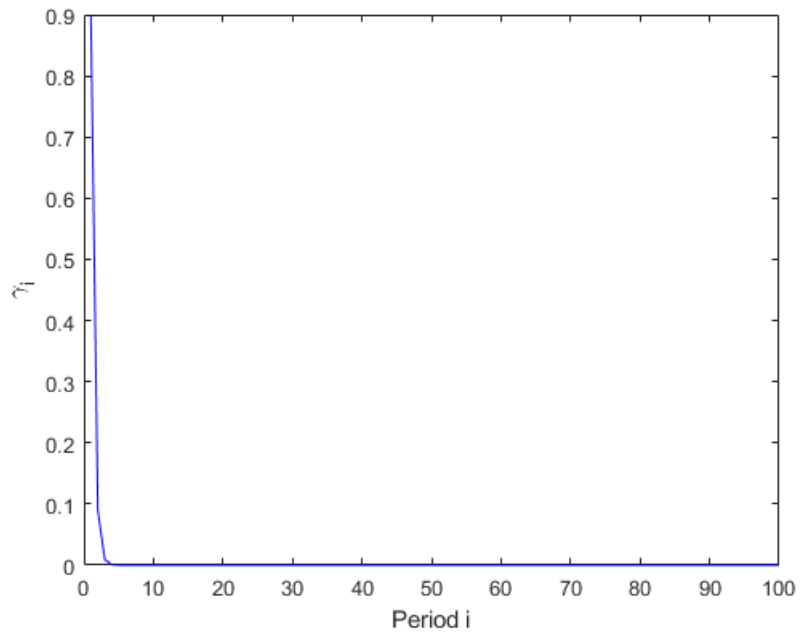


Figure 2.1: The equilibrium allocation of stakes when the agent set stakes. $\underline{\delta} = 0.1$ and $n = 100$.

Commitment bad agent

The first thing to realize is that any δ_i for which the (good) agent strictly prefers action 1 can not be an equilibrium allocation. The reasoning is that the social planner's objective function is increasing in δ_i , and he can profit by deviating to a higher δ_i . Moreover, any stake allocation where the agent strictly prefers action 0 can also not be an equilibrium strategy. Following the analysis in Chapter 1 1.3.1, the agent is not hired in period i if $\mu_i > \mu'_i$, where $\mu'_i < 1$. Thus, the social planner sets the stakes so the agent is indifferent between both actions in any period. Note that this holds in any history where no action 0 is observed in the previous periods. If a history h_i involves action 0 for some periods, then $\rho_i = 1$. After any such history, the agent is hired for the remainder of the game. The choice of δ_i for such periods is arbitrary.

I define ranges for the reputation of the agent as in chapter 1 Section 1.4. Define $P(t) = \frac{(1-2k)2^{t-2}}{(1-2k)2^{t-2} + k}$ where n denotes the total number of periods and $t \in \{2, 3, \dots, n\}$. Note that for $\rho_1 \geq P(n)$, the agent does not play action 0 with a positive probability in the equilibrium. Moreover, the agent is never hired for $\rho_1 < P(1)$. Thus, I focus on $t \in \{1, 2, 3, \dots, n\}$ again.

Proposition 2.2. *Let $P(t) \leq \rho_1 < P(t+1)$, the following assessment constitutes the agent-optimal perfect Bayesian equilibrium.*

Hiring starts in the first period. After histories in which action 0 is never observed, $\delta_i = \frac{2^{t-i}(1+2k) - 2k}{2^{t-i}(1+2k) - k}$ in period i . The agent is hired for a total of t periods if no action 0 is observed. $\mu_i \in [0, \mu'_i]$ after such histories.

The agent will be hired until the end of the game after any history in which action 0 is observed. $\mu_i = 0$ and $\delta_i \in [\underline{\delta}, \bar{\delta}]$ after such histories.

Proof. Appendix B. □

In histories where no action 0 is observed, the stakes evolve as follows:

$$\Gamma = \left\{ \frac{k}{2^{t-1}(1+2k) - k}, \frac{2k}{2^{t-1}(1+2k) - k}, \dots, \frac{2^{t-1}k}{2^{t-1}(1+2k) - k} \right\}$$

The interaction starts small and gradually increases by multiplying by 2. Note that the stakes do not add up to 1. The reasoning is that in period t , $\frac{2^{t-1}}{2^{t-1}(1+2k) - k}$ proportion of the total stakes is left to the future¹. If no action 1 is observed in any period, these stakes are lost.

Strategic bad agent

The social planner sets the stakes so that the bad agent is indifferent between both actions in any period. In the social-planner-optimal equilibrium, $\mu_i = 0 = \nu_i$ in all periods. The reasoning is that the social planner's continuation payoff decreases in ν_i . If a history h_i involves action 0 for some periods, then $\rho_i = 1$. After any such history, the agent is hired for the remainder of the game. The choice of delta for such periods is arbitrary.

Proposition 2.3. *Let $P(t) \leq \rho_1 < P(t+1)$, the following assessment constitutes the agent-optimal perfect Bayesian equilibrium.*

Hiring starts in the first period. After histories in which action 0 is never observed, $\delta_i = \frac{1 + (t-i)k}{1 + (t+1-i)k}$ in period i . The agent is hired for a total of t periods if no action 0 is observed. $\mu_i = 0 = \nu_i$ for all $i \in N$.

The agent will be hired until the end of the game after any history in which action 0 is observed. $\delta_i \in [\underline{\delta}, \bar{\delta}]$ after such histories.

Proof. Appendix B. □

In the histories where no action 0 is observed, all periods receive the same stake of $\gamma_i = \frac{k}{1 + tk}$. Note that this is not the same as the original analysis, where all the periods receive the same stake. This allocation is valid for the periods where no action 0 is observed in history. This also implies that if no action 0 is observed until the period t , then $\frac{1}{1 + tk}$ proportion of all the stakes available is wasted. Note

¹It is attained by $\prod_1^t \delta_i$

that γ_i is inversely related to t , i.e., higher t implies lower stakes. It may be surprising because the higher t , the higher the initial reputation. One may expect the initial stakes to be lower for the lower reputation levels. However, the reasoning is the reverse. If the agent has a lower initial reputation, she holds higher reputational concerns. Hence, to make the bad agent indifferent between investing in the reputation or not, the interaction has to start higher.

2.5 Long-term principal

2.5.1 Agent sets stakes

This section will analyze the main model with a long-term principal where the agent sets relative stakes of the interaction. I show that if the reputation is not so low, the agent allocates stakes so that (1) the principal hires the agent and (2) the agent has no reputation incentive. This allocation enables the agent to maximize her continuation payoff. I start with the bad commitment type analysis.

Commitment bad agent

Consider any critical period where the reputation is sufficiently small so that the agent will not be hired after action 1. The agents prefer to leave the least possible stakes to the future. However, if the current period receives high stakes and the agent's reputation is low, the principal does not hire the agent. To avoid this, the agent increases δ_i , and if the reputation is sufficiently low, then the agent's expected payoff maximizing stake allocation involves starting the period small. On the other hand, if the agent's reputation is not sufficiently low in period i , the agent sets stakes as low as possible in any period she is hired.

The following proposition formalizes the equilibrium.

Proposition 2.4. *If the agent chooses the stakes, then in the agent-optimal perfect Bayesian equilibrium the agent is hired as long as $\rho_i \geq \frac{\underline{\delta}(1-2k)}{2(1-\underline{\delta})k}$. The agent's equilibrium strategy is as following:*

1. If $\rho_1 \geq P(n)$: $\mu_i = 0$ and $\delta_i \in [\underline{\delta}, \bar{\delta}]$, $\lambda_i = 1 \forall i \in N$.

2. If $\frac{(k - \underline{\delta})(1 - 2k)}{(k - \underline{\delta} + \underline{\delta}k)} \leq \rho_1 < P(n)$: $\mu_1 = 0$ and $\delta_1 = \max \left\{ \underline{\delta}, \frac{1 - 2k - \rho_1}{1 - 2k - \rho_1 + \rho_1 k} \right\}$.

If $a_j = 0$ is observed for some $j \in N$, then $\mu_i = 0$ and $\delta_i \in [\underline{\delta}, \bar{\delta}]$, $\lambda_i = 1 \forall i \in \{j, j + 1, \dots, n\}$. Otherwise, if no action 0 is observed before period $i \in N/1$:

(a) If $\frac{(k - \underline{\delta})(1 - 2k)}{(k - \underline{\delta} + \underline{\delta}k)} \leq \rho_i < \frac{(1 - 2k)2^{n-1-i}}{(1 - 2k)2^{n-1-i} + k}$: $\mu_i = 0$ and

$$\delta_i = \max \left\{ \underline{\delta}, \frac{1 - 2k - \rho_i}{1 - 2k - \rho_i + \rho_i k} \right\}.$$

(b) If $\frac{\underline{\delta}(1 - 2k)}{2(1 - \underline{\delta})k} \leq \rho_i < \frac{(k - \underline{\delta})(1 - 2k)}{(k - \underline{\delta} + \underline{\delta}k)}$: $\mu_i = 1$ and $\delta_i = \bar{\delta}$.

3. If $\frac{\underline{\delta}(1 - 2k)}{2(1 - \underline{\delta})k} \leq \rho_1 < \frac{(k - \underline{\delta})(1 - 2k)}{(k - \underline{\delta} + \underline{\delta}k)}$: $\mu_1 = 1$ and $\delta_1 = \bar{\delta}$. The agent fully screens her type and hired for the remainder of the game. $\mu_i = 0$ and $\delta_i \in [\underline{\delta}, \bar{\delta}]$, $\lambda_i = 1 \forall i \in N/1$.

Proof. Appendix B. □

If the initial reputation of the agent is not sufficiently small, then she starts the interaction large. The relative stakes gradually decrease as long as no action 0 is observed in history. If, in any period, the agent's reputation falls below $\frac{\underline{\delta}(1 - 2k)}{2(1 - \underline{\delta})k}$, the agent is not hired if she puts high stakes to period i . In such period i , the agent needs to reveal her type. She sets high enough δ_i so that it is rational to play $\mu_i \geq 0$. Moreover, if $\mu_i \geq 0$, her expected continuation payoff is increasing in δ_i . Thus, $\delta_i = \bar{\delta}$.

I know the equilibrium behavior if the initial reputation of the agent is sufficiently small. She will choose to start small to take the costly reputational action in the first period. The higher δ_1 , the lower the cost in the first period is. Hence, the agent

sets $\delta_1 = \bar{\delta}$, screens her type, and is hired for the remainder of the game. Note that, for $\bar{\delta}$ close enough to 1, the agent is hired for the initial reputation levels close to 0. Hence, endogenizing the allocation of relative stakes makes the interaction possible even if there is a very small probability of the agent being the good type.

Strategic Bad Agent

A critical difference in the equilibrium behavior when I introduce a strategic bad agent is that the good agent no more starts small for any initial reputation levels. The rough intuition is as follows. If the bad agent is a commitment type, the good agent can leave as many important decisions as possible to the future and perfectly screen her type by playing action 0. However, if the bad agent is also a strategic player, then it is not feasible. If a sufficiently high proportion of the decisions are left to the future, then the bad agent also plays action 0 with a positive probability. It slows down the reputation evaluation, which results in a lesser expected continuation payoff to the good agent. Hence, this behavior becomes less appealing to the good agent, and she prefers to start large as long as she believes the principal will hire her. Recall:

$$R(t, n) = \frac{(1 - 2k)2^{t-1}}{(1 - 2k)2^{t-1} + (n - t + 2)k}$$

I defined $R(t, n)$ to set t corresponding to initial reputation level. Now, I will define general $R(i, t, n)$ for period i :

$$R(i, t, n) = \frac{(1 - 2k)2^{t-1}}{(1 - 2k)2^{t-1} + (n - t - i + 3)k}$$

so that if $\rho_i > R(i, t, n)$ and the agent has no reputation concern in periods $i, i + 1, \dots, i + t - 2$, then in the worst history in which no action 0 is observed $\rho_{i+t-1} \geq 1 - 2k$. The agent's reputation stays above $1 - 2k$ for at least t periods. Hence, if $\rho_i \geq R(i, n - i + 1, n)$, then the agent knows that if she does not take the costly reputation action, then $\rho_n \geq 1 - 2k$ and she will be hired for all the remaining periods. If that is the case then the agent's choice of δ is arbitrary and the agent plays stage-game maximizing strategy.

The agent's choice of δ_i depends on ρ_i as follows. For $\rho_i \geq 1 - 2k$, the principal hires the agent irrespective of the allocation of stakes, because the principal obtains a non-negative payoff in period i (if there is no reputation concern in period i , $U_i^P = k - \frac{1}{2} + \frac{1}{2}\rho_i$). Consider the extreme case where $\underline{\delta} = 0$ and the agent sets $\delta_i = \underline{\delta}$. The interaction becomes one-shot as there are no stakes left to the future. From the analysis before, the principal hires the agent if $\rho_i \geq 1 - 2k$. In that case, it is optimal for the agent to set δ as low as possible, i.e. leave the minimum possible stakes to the future.

On the other hand, if the reputation is below $1 - 2k$, then the minimum $\delta_i = \underline{\delta}$ leads to negative continuation payoff to the principal (by definition, $\underline{\delta}$ is sufficiently small). Hence, the agent sets minimum δ_i value which leads to a non-negative continuation payoff for the principal, which corresponds to $\delta'_i = \frac{1 - 2k - \rho_i}{1 - 2k - \rho_i + \rho_i k}$. Again, the agent chooses the minimum possible δ value. For $\rho_i \geq \frac{1}{2} - k$, $\delta'_i \leq \frac{1}{1 + k}$. Thus, $\mu_i = 0$ and $\nu_i = 0$ is an equilibrium strategy.

Consider $\rho_i < \frac{1}{2} - k$. The lowest δ_i leading to a non-negative payoff for the principal is greater than $\frac{1}{1 + k}$. Therefore, no reputation concern is not an equilibrium strategy implied by the sequential rationality of the agent. Both types of agents would deviate to play action 0. This implies that, given $\rho_i < \frac{1}{2} - k$, the agent may only be hired in period i if the agent has reputation concerns. Suppose the principal hires the agent and the agent has reputation concerns in period i . Consider the equilibrium path following period i . The agent's payoff in period $i + 1$ is decreasing in $\underline{\delta}$. Define $\tilde{\delta}_i$ as agent's choice of δ_i which implements reputation concerns in the equilibrium. As V_{i+1}^g is decreasing in $\underline{\delta}$; lower $\underline{\delta}$ implies lower $\tilde{\delta}_i$. To see the intuition, consider the following. If δ_{i+1} increases, the agent's payoff decreases in period $i + 1$. From the period i agent's point, the future becomes less appealing. Hence, to have reputation concerns, δ_i should be higher ($\tilde{\delta}_i$ increases). Next thing to realize is that the principal gets a negative payoff in period i as there are reputation incentives.

Hence, for the principal to hire the agent, equilibrium δ_i needs to be above a certain level. In words, the higher δ_i , the lower stakes are attached to period i in which the principal obtains a negative payoff. Returning to the sequential rationality constraint of the agent; for indifference δ_i to be high, $\underline{\delta}$ should be above a certain level.

Overall, for the principal to hire the agent in period i where $\rho_i < \frac{1}{2} - k$ (and the agent has reputation incentives), equilibrium δ_i should be above a certain level (which implies $\underline{\delta}$ should be above a certain level). However, as I focus on the sufficiently small $\underline{\delta}$, I conclude that the principal does not hire the agent in such a case.

Remark 2.1. *If $n < \frac{1}{k} + 1$, the agent has no reputation concern in the equal-stakes model. In such case, $\frac{1}{2} - k < R(1, n)$. It implies the agent is hired for lower reputation level if she sets the stakes. On the other hand, consider the special case in*

Proposition 1.7. The agent is hired if $\rho_1 \geq \frac{(1-2k)2^{n-3}}{(n-1)k(2^{n-2}-1)}$. A simple calculation

shows that $\frac{(1-2k)2^{n-3}}{(n-1)k(2^{n-2}-1)} < \frac{1}{2} - k$. It implies that the agent is hired for lower reputation levels if the stakes are equally allocated. Hence, the agent's ability to set stakes hurts her equilibrium payoff.

In any period i , the agent is hired if and only if $\rho_i \geq \frac{1}{2} - k$. Define

$$P(1, t, n) = \begin{cases} \frac{(1-2k)2^{t-1}}{(1-2k)2^{t-1} + 1 + 2k}, & t \in \{1, \dots, n-2\} \\ R(1, t, n), & t = n-1 \end{cases}$$

If $P(1, t, n) \leq \rho_1 < P(1, t+1, n)$, $\mu_i = \nu_i = 0$ and $a_i = 1$ for $i \in 1, 2, \dots, t-1$, then

$\rho_t \geq \frac{1}{2} - k$ and $\rho_{t+1}(1) < \frac{1}{2} - k$. In words, if the agent does not hold reputation concerns, then she is hired for exactly t periods if no action 0 is observed. The

principal hires the agent in period $i < n$ if $\rho_i \geq \frac{1}{2} - k$. Hence, given $P(1, t, n) \leq \rho_1 < P(1, t+1, n)$, the agent is hired for t periods if no action 0 is observed. The

range of reputation is defined so that, $\rho_t \geq \frac{1}{2} - k$ and $\rho_{t+1} < \frac{1}{2} - k$ if no action 0 is observed for $t < n-1$. This implies, the agent is hired in period t but not in

period $t + 1$ onward. Note that $P(1, 1, n)$ corresponds to $\frac{1}{2} - k$, hence, the agent is not hired if $\rho_1 < P(1, 1, n)$.

However, in period n , the principal hires the agent if the reputation is not below $1 - 2k$, because there are no more periods left to be played. If $P(1, n - 2, n) \leq \rho_1 < P(1, n - 1, n)$, the agent is hired for $n - 1$ periods if no action 0 is observed. Hence, $\rho_{n-1} \geq \frac{1}{2} - k$ and $\rho_n < 1 - 2k$. Thus, the agent is hired in period $n - 1$, but not hired in period n . As there is no reputation concern, $\rho_{i+1}(0) = 1$ and $\rho_{i+1}(1) = \frac{\rho_i}{2 - \rho_i}$ for all i .

If $\rho_1 \geq P(1, n, n)$, the agent has no reputation concerns and is hired in all periods irrespective of the observed action. By Bayes' rule, $\rho_n \geq 1 - 2k$ even if no action 0 is observed. Both types of agents get the maximum possible payoff - k - from the project. Hence, this equilibrium is not payoff-dominated. The agent arbitrarily sets δ_i as she is hired every period.

Consider the following proposition.

Proposition 2.5. *Agent-optimal perfect Bayesian equilibrium is as following. $\mu_i = 0 = \nu_i$ and*

$$\delta_i \in \begin{cases} [\underline{\delta}, \bar{\delta}], & \rho_i \geq R(i, n - i + 1, n) \\ \underline{\delta}, & 1 - 2k \leq \rho_i < R(i, n - i + 1, n) \\ \delta'_i = \frac{1 - 2k - \rho_i}{1 - 2k - \rho_i + \rho_i k}, & \frac{1}{2} - k \leq \rho_i < 1 - 2k \end{cases}$$

$$\lambda_i = \begin{cases} 0, & \rho_i < \frac{1}{2} - k \\ 1, & \rho_i \geq \frac{1}{2} - k \end{cases}$$

Proof. Appendix B. □

The stakes start large and gradually decrease. For $P(1, t, n) \leq \rho_1 < P(1, t + 1, n)$, the agent is hired for at least t periods. After s periods, the reputation of agent falls below $1 - 2k$ if no action 0 is observed. Note that $s < t$ as the agent is also hired

when $\rho_i \leq 1 - 2k$. The evolution of stakes is as follows:

$$\Gamma = \{(1 - \underline{\delta}), (1 - \underline{\delta}) \underline{\delta}, (1 - \underline{\delta}) \underline{\delta}^2, \dots, (1 - \underline{\delta}) \underline{\delta}^{t-3}, \underline{\delta}^{t-2} (1 - \delta'_{t-1}), \underline{\delta}^{t-2} \delta'_{t-1} (1 - \delta'_t)\}$$

The agent may survive for more than two periods after the reputation falls below $1 - 2k$. It solely depends on the k parameter. In the distribution above, it was implicitly assumed that $k > 1/6$, so that if $\rho_{t-1} < 1 - 2k$ and action 1 is observed for two consecutive periods (periods $t - 1, t$), then $\rho_{t+1} < 1/2 - k$. For any k parameter, the results of starting large and gradual decrease hold.

2.5.2 The long-term principal sets relative stakes

This section analyzes the framework where the long-term principal sets the relative stakes. The principal always allocates the stakes, giving the agent a reputational incentive. The difference stems from the difference in continuation payoffs. If δ_i is such that no type has a reputational incentive, then the agent's expected continuation payoff decreases in δ_i . The reasoning is that the agent gets the payoff of k in period i , whereas the future expected continuation payoff is strictly less than k . Consider the principal. For such δ_i , as there will be the complete revelation of the type of the good agent after action 0, the future expected continuation payoff of the principal is higher than the payoff in period i . Thus, the principal prefers to increase δ_i until the point where the reputational incentives emerge.

Note that $\delta_i < \frac{1}{1+k}$ is the sufficient condition for the absence of reputation concerns. In other words, if period i receives more than $\frac{k}{1+k}$ proportion of the remaining decisions, then the agent has no reputation concern in any equilibrium. It follows from the sequential rationality of the agent. In any period i , the maximum payoff agent can get from taking the costly reputation-driven action is given by $(1 - \delta_i)(k - 1) + \delta_i(k)$. On the other hand, she can get a payoff weakly greater than $(1 - \delta_i)k$ if she plays action 1. Hence, if $\delta_i < \frac{1}{1+k}$, then the agent plays stage-game maximizing action.

Proposition 2.6. For $\rho_1 < \frac{2^{n-2}(1-2k)}{2^{n-2}(1-2k)+k}$, there is no equilibrium where $\lambda_1 = 1$ and $\delta_1 < \frac{1}{1+k}$.

Proof. Appendix B. □

Remark 2.2. Consider period 1 with endogenous stakes, $\gamma_1 = 1 - \delta_1$. $\delta_1 < \frac{1}{1+k}$ implies that $\gamma_1 > \frac{k}{1+k}$. Recall the sufficient condition for the absence of reputation concerns in the benchmark model stated in Proposition 1.1. I showed that if $n < \frac{1}{\frac{1}{k} + 1}$, then the agent has no reputation concern in any equilibrium. Note that, in benchmark model, each period receives a stake of $\frac{1}{n}$, i.e. $\gamma_i = \frac{1}{n} \forall i$. The condition $\delta_1 < \frac{1}{1+k}$ corresponds to $n < \frac{1}{\frac{1}{k} + 1}$ for $\gamma_1 = \frac{1}{n}$.

The principal setting stakes enables him to hire the agent for low initial reputation levels. The reasoning is that as the principal can manage the reputation concerns of the agent, he can strategically allocate stakes to keep the agent's reputation concerns at the desired level.

Commitment bad agent

I define the principal-optimal equilibrium, where the bad agent is a commitment type. Suppose that the good agent is a strategic type who maximizes her expected continuation payoff while the bad agent is a commitment type. Then it is optimal for the principal to put high stakes to the future interaction so that the good agent screens her type. The reasoning is straightforward. In any period i where $\rho_i < 1$, the principal gets a negative payoff with probability $(1 - \rho_i)$. If $\bar{\delta}$ is sufficiently high, then the principal sets $\delta_i = \bar{\delta}$. The minimum stake is attached to period i . The good agent screens her type with a probability of $1/2$ in each period. If no action

0 is observed for a sufficiently long time, then the agent's reputation approaches 0. In such a period where the agent will no more survive after action 1, she has reputational concerns and plays action 0 with probability 1 and screens her type.

The principal's expected continuation payoff approaches to $\rho_1 k$ in the equilibrium as $\bar{\delta}$ approaches to 1. This is the maximum continuation payoff the principal can get and can only be attained in a framework with no informational asymmetry, i.e., the principal observes the agent's type.

Proposition 2.7. *In the agent-optimal perfect Bayesian equilibrium where $\bar{\delta}$ is sufficiently high, $\delta_i = \bar{\delta}$ for all $i \in N$. $\mu_i = 0$ for all i except when ρ_i close enough to 0 and $\mu_i = 1$.*

Proof. Appendix B. □

By definition, $\delta_i = \bar{\delta}$ implies $\gamma_i = (1 - \bar{\delta}) \bar{\delta}^{i-1}$ for all $i \in N/n$ and $\gamma_n = \bar{\delta}^{n-1}$. Fixing $\bar{\delta} = 0.9$ and $n = 100$, Figure 2.2 represents the evolution of the stakes in the equilibrium. Observe that the interaction starts small in the sense that maximum possible stakes are left to the future in each period. Stakes slowly decrease initially until the last period.

Strategic bad agent

This section will define the principal-optimal equilibrium for the sufficiently small initial reputation. If the initial reputation is low, the principal starts the interaction with low stakes. I define a sufficiently small initial reputation as the lowest initial reputation level for which the agent is hired in the equilibrium. By doing so, the principal gradually updates his belief about the agent's type. In the equilibrium, the whole interaction (except for period n) serves as the *evaluation phase*.

I define *evaluation phase* as follows: in any period of the evaluation phase, the principal continues to hire the agent as long as action 0 is observed; if otherwise, action 1 is observed in any period, the principal ends the relationship. $\nu_i > 0$ must hold in the evaluation phase except for the last period. If $\nu_i = 0$ for some i , then $\rho_{i+1}(0) = 1$ which means the agent is hired for the remaining periods if $a_i = 0$ is

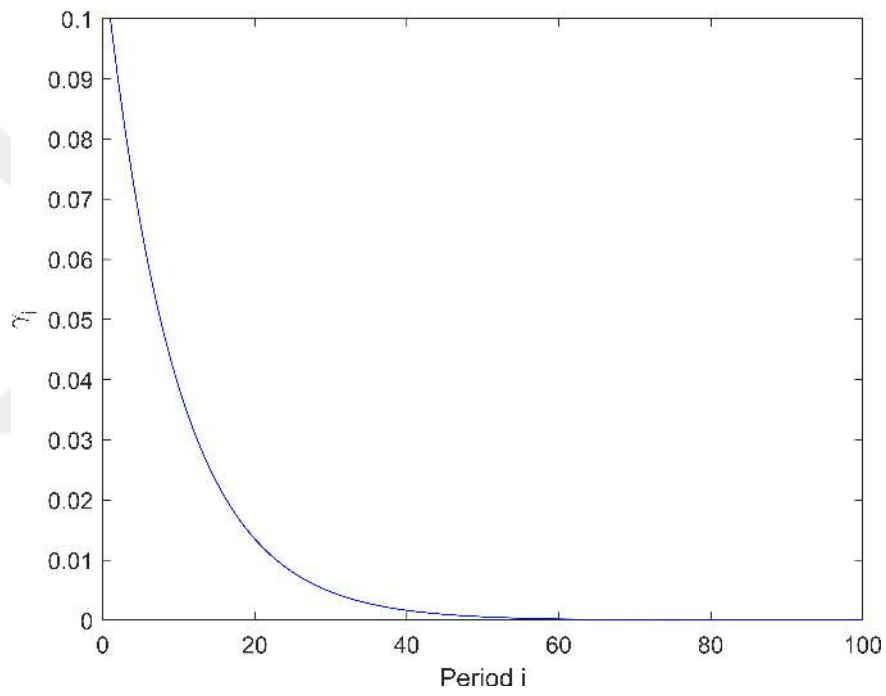


Figure 2.2: The equilibrium allocation of stakes when the long-term principal set stakes and the bad agent is a commitment type. $\underline{\delta} = 0.1$ and $n = 100$.

observed. It contradicts the definition of the evaluation phase. $\nu_i = 0$ can hold only in the last period of the evaluation phase.

Consider the following cases to understand how the principal strategically sets the evaluation phase's length. Case 1: The evaluation phase involves the first s periods, i.e., given no action 1 is observed in the previous periods, the agent is hired for periods $s + 1, s + 2, \dots, n$ irrespective of the action observed. Case 2: Evaluation phase involves the first t periods such that $s < t$. In case 1, more periods are left to be played after the evaluation phase. Thus, if each period received the same stake, the agent would have more reputation concerns in case 1. Hence, the principal needs to put higher stakes per the period of the evaluation phase in case 1 compared to case 2. Thus, the smaller the evaluation phase is, the higher the principal risks in each period. The intuition is that the shorter the evaluation phase is, the higher stake each period should receive to incentivize the agent. However, it means risking high stakes, which leads to a lower expected continuation payoff. If the initial reputation is sufficiently small, the evaluation phase involves all periods except the last one.

Building reputation is costlier for the bad agent except for the period $n - 1$. In period $n - 1$, both types' future expected payoffs are the same; hence, they have the same sequential rationality constraint. For earlier periods, playing action 0 costs the same to both types, whereas the good agent has a higher future expected payoff. By taking the costly action, the good agent gets more chances to improve her reputation costlessly in the future (when the true state is 0 and she plays accordingly). So the good agent values future interaction more. It implies that investing in her reputation is less costly for her. Thus, the sequential rationality constraints of both types differ when there is more than one period left to be played. In other words, if the stakes are allocated so that $\nu_i > 0$, then $\mu_i = 1 \forall i \in N/\{n - 1, n\}$.

The principal sets δ_i values so that the bad agent is indifferent between two actions in each period and $\nu_i \in (0, 1) \forall i \in N/\{n - 1, n\}$. If $a_i = 1$ is observed in any period, the principal is sure that the agent is the bad type. Whereas the reputation after $a_i = 0$ evolves at just enough pace so that the agent is hired for one more period.

Remark 2.3. In the equilibrium, there exists $\underline{\nu}_i$ and $\bar{\nu}_i$ for all $i \in N$ such that $\nu_i \in [\underline{\nu}_i, \bar{\nu}_i]$ is consistent with equilibrium beliefs. For $\nu_i \in [\underline{\nu}_i, \bar{\nu}_i]$, $\rho''_{i+1} \geq \rho_{i+1}(0) > \rho'_{i+1}$, in words, reputation stays in the lowest range. As the principal's payoff decreases in ν_i for all i , the bad agent plays action 0 with the lowest possible probability that satisfies equilibrium conditions in the principal-optimal equilibrium. Thus, $\nu_i = \underline{\nu}_i$ and the reputation evolves as $\rho_{i+1}(0) = \rho''_{i+1} \forall i \in N$. The reputation evolves in just enough pace so that $\rho_{n-1}(0) = \frac{1-2k}{1-k}$.

Proposition 2.8 formalizes the equilibrium behavior.

Proposition 2.8. There exists $\rho'_1 < \rho''_1$ such that for $\rho'_1 \leq \rho_1 < \rho''_1$, the following strategies constitute the principal-optimal perfect Bayesian equilibrium:

$$\nu_i \in (0, 1), \mu_i = 1 \text{ for } i \in N/\{n, n-1\}, \nu_{n-1} = 0, \mu_{n-1} \in [0, 1], \nu_n = 0, \mu_n = 0$$

$$\delta_i^* = \frac{\sum_{h=0}^{n-i-1} k^h}{\sum_{h=0}^{n-i} k^h}, \lambda_i(0) = 1 \text{ and } \lambda_i(1) = 0 \text{ for all } i \in N/1. \lambda_1 = 1$$

Proof. Appendix B. □

In period $n-1$, $\delta_i = \frac{\sum_{h=0}^{n-i-1} k^h}{\sum_{h=0}^{n-i} k^h} = \frac{1}{1+k}$ and the good agent is also indifferent between two actions and she mixes. The sequential rationality constraints of both types are the same in $n-1$. There is one-period play left, so the future expected payoff of both types is k if the game is not terminated and 0 if it is terminated. As the principal's continuation payoff is decreasing in ν_i , then $\nu_{n-1} = 0$ in the principal-optimal equilibrium. I see that if I start with small enough initial reputation, principal starts with low γ and increases it geometrically through interaction, which is in accordance with the “starting small” literature - $\gamma_i = \frac{k^{n-i}}{\sum_{h=0}^{n-1} k^h}$. From the equilibrium analysis above, the principal does not prefer the agent to screen her type before the period $n-1$. The intuition is that the principal needs to put high stakes in the evaluation periods to incentivize the agent. But it means risking high stakes, which leads to a lower expected continuation payoff.

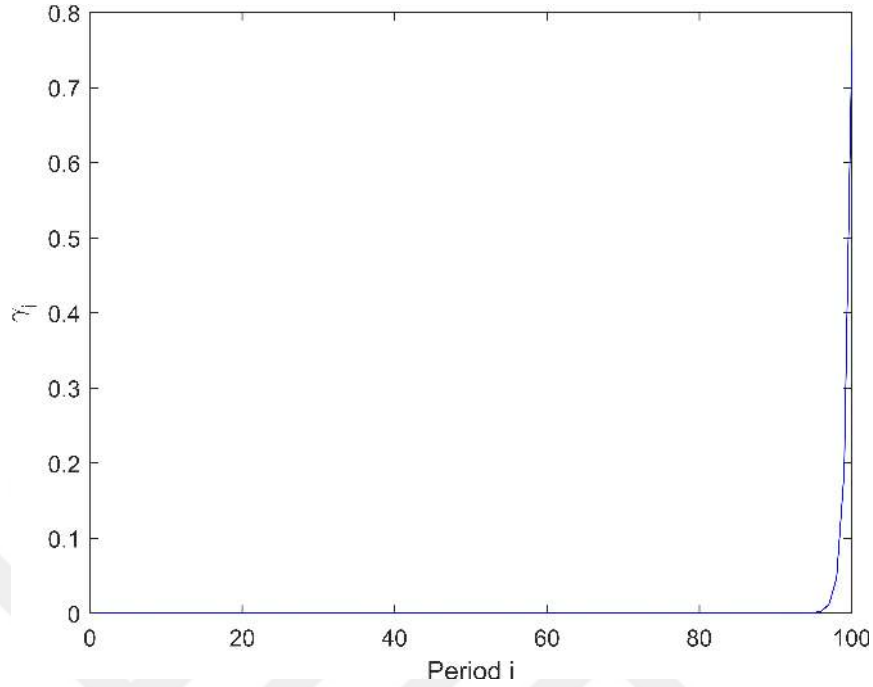


Figure 2.3: The equilibrium allocation of stakes when the long-term principal set stakes and the bad agent is a strategic type. $k = \frac{1}{4}$ and $n = 100$.

Fixing $k = \frac{1}{4}$ and $n = 100$, Figure 2.3 represents the evolution of the stakes in the equilibrium. Observe that the interaction starts very small and gradually increases until the last period.

2.6 Conclusion

I analyzed a repeated interaction between a principal and an agent involving different allocations of relative stakes. I borrowed the model from Ely and Välimäki [2003], started with providing a finite version of the model offered by EV, and showed that as the number of periods converges to infinity, equilibrium behavior converges to EV.

The main contribution of this chapter is introducing endogenous stakes. I showed that the player's optimal stake allocation depends on their position on the asym-

metric information problem. The agent is informed, and there is no uncertainty about the principal's preferences. Hence, the agent prefers to start the interaction large so there are no reputation concerns. Starting large may or may not be to the advantage of the agent. If the agent does not have reputation concerns in the equal-stakes allocation, then starting large has a two-fold advantage for the agent. First, the agent gets a higher payoff in starting-large allocation. Second, the agent is hired for lower initial reputation levels for which she is not hired if the stakes are equally allocated.

Nevertheless, the agent's ability to endogenously set the stakes may affect negatively. I show the existence of cases in which the agent is hired if the stakes are equally allocated but not hired if the agent sets the stakes. The intuition is as follows. If the agent starts large, it is risky for the principal to hire the agent if the initial reputation is sufficiently low. On the other hand, suppose the agent starts with small stakes. Then there will be reputation concerns. The agent can not commit to an allocation of stakes at the beginning of the interaction. Screening occurs in the initial period(s), and the principal continues hiring the agent if the reputation exceeds a certain level. However, after the reputation level increases, the agent puts the highest possible stakes in every consecutive period. The principal obtains a zero or sufficiently small continuation payoff in any such period. In addition, the principal gets a negative period payoff in any period involving reputation concerns. Thus, the principal expects a negative continuation payoff in the first period if the agent has a sufficiently small initial reputation level. In such cases, exogenously set stakes may be preferred to endogenous stakes.

The principal does not know the agent's preferences and is uninformed about the state of the world. For him, starting large is equivalent to starting a risky business with high stakes. On the other hand, starting small is appealing to the principal for the following reasons:

1. As larger stakes are left to the future, the agent will hold reputation incentives.

The principal benefits from the reputation incentives of the agent as it helps him update his beliefs on the agent's type.

2. As a smaller stake is allocated to the current period, the principal risks a small proportion of the project. If the agent is a bad type, then her harm to the project will be limited.

Hence, the principal prefers to start small and always benefits from the ability to set the stakes endogenously.



Chapter 3

GRADUALISM IN A PRINCIPAL AGENT INTERACTION - EXPERIMENT

3.1 Introduction

We frequently participate in interactions that last longer than a single period. In environments with asymmetric information, repeated play becomes valuable. Reputation incentives arise, and the players can obtain higher equilibrium payoffs via updating beliefs in each other. Furthermore, repeated play enhances cooperation in prisoners' dilemma-like settings where it is hard to reach cooperation in one-time interactions.

This chapter theoretically and experimentally analyzes a repeated principal-agent interaction. The principal decides whether to delegate a job to the agent whose preferences are unknown. I focus on the setting where there is a sufficiently small probability of the agent being a good type. If the interaction is one-shot, the principal does not hire the agent with such a low reputation. However, in a repeated play, the principal can monitor the agent's behavior and act accordingly. The principal hires the agent in the first period and continues hiring as long as the agent displays the desired behavior. The principal's hiring behavior creates reputational incentives for the agent to act according to the principal's expectations.

The extent of the reputational incentives depends on how the stakes are allocated across the interaction periods. In the classical studies of repeated communications, each period receives an equal weight. In this chapter, I vary the allocation of stakes to find the principal-optimal one. For instance, the agent does not have reputational incentives if the interaction starts with sufficiently high stakes. In other words, if the future communication is not adequately important compared to the current

play, then the agent aims to maximize today's payoff rather than focusing on the reputation. It may be the desired behavior in specific environments. However, the principal does not hire the agent in the absence of reputational incentives if her initial reputation is low. The principal-optimal allocation involves starting small and increasing the stakes gradually.

This chapter analyzes a three-period repeated principal-agent game. The agent is either a good or a bad type, and the principal does not know the agent's type. The principal hires the informed agent to observe the state and take action accordingly. The good agent shares the same preferences with the principal and prefers to take the action that matches the state. On the other hand, the bad agent prefers a particular action irrespective of the state. Consider, for example, a state official hiring a private firm for a public project that will last up to three years. At the beginning of each year, the official decides whether to hire the agent or not. If the agent is hired, she implements the project. At the end of the year, the official observes the firm's action and whether the execution was adequate. Next, he makes the hiring decision for the following year.

It is optimal for the principal to start the interaction small in such a setting. Starting small leads to reputation incentives that benefit the principal. In the initial periods, the bad agent does not want to reveal her type and displays the principal's desired behavior. Morris [2001] names this as the "discipline effect" of reputation incentives. Following the example above, the government official may be free to choose the project's relative size each year. He can dynamically decide which proportion of the remaining project to implement in each period. This way, he can control the reputational incentives of the agent. Alternatively, a third party may allocate the stakes across the periods. The optimal allocation of stakes is such that the bad agent is indifferent between investing in reputation or not in each period.

In the experiment performed in this chapter, I test the theory's predictions with four treatments. Each treatment involves a particular allocation of the stakes. In one treatment, each period gets an equal stake. There are three treatments for starting small. Stakes grow at a different speed in each treatment. I show that the

smaller the interaction starts, the higher the reputation incentives are. The reputation incentives of the agent benefit the principal to some extent. The principal's payoff is higher in the starting small treatments compared to the equal-stakes treatment. Hence, if managed, these incentives can improve the equilibrium payoff of the principal. Reputation incentives become more valuable in a setting, as presented in this chapter, where there is a high risk of hiring the agent.

3.2 Literature Review

My analysis lies in the intersection of different branches of the literature. First, it belongs to the literature on repeated interactions. Papers by Kreps and Wilson [1982] and Milgrom and Roberts [1982b] were the first to formalize the effects of the reputation that pioneered the conventional wisdom that the reputation concern increases commitment power. They showed that introducing repeated actions and information asymmetry leads to a reputation effect. I borrow the theoretical foundation from Chapter 1. The model provided in this chapter departs from Chapter 1 in two aspects. First, the principal never observes the state of the world in the initial chapters. In this chapter, the principal observes the agent's action and the true state of the world. Secondly, I exogenously set the interaction's relative stakes for simplicity during the experiments. We also share the same stage game with Ely and Välimäki [2003]. They analyze an infinitely-repeated principal-agent game. The interaction in my model has a finite horizon. In addition, I have varying relative stakes in my analysis.

Second, there is a growing theoretical literature on the effect of starting small in various repeated interactions. Most notable theoretical works include Sobel [1985], Ghosh and Ray [1996], Kranton [1996], Watson [1999, 2002b], Atakan et al. [2020]. My model separates from earlier works on the stage game. I contribute to the literature by showing that starting small is a valuable tool for improving payoffs in a different environment.

Third, there is a substantial amount of experimental work analyzing repeated

interactions.¹ The closest work to mine is by Grosskopf and Sarin [2010], who provided experimental work where a long-run player faced short-run players. They found that the reputation is seldom harmful, and its beneficiary effects are not tremendous as thought. They analyzed a repeated lender-borrower setup and concluded that reputation incentives do not substantially affect the players' behavior. More importantly, other-regarding preferences seem to triumph the reputation concerns. I reach the opposite result. I make the distinction between other-regarding preferences and reputation concerns as follows. The bad agent prefers the blue color in the one-shot game. Hence, if we observe the bad agent choosing green when the state of the world is green, there are two possible explanations. First, it may stem from other-regarding preferences. Consider the following scenario. The state of the world happens to be green for two consecutive periods. In the first one, the bad agent chooses blue - maximizes his stage game payoff. However, this action hurts the principal. Then, in the second period, he decides to choose green. This way, she gets a lower payoff but achieves a more fair allocation of payoffs between herself and the principal. The second explanation is the reputation investment. The bad agent may choose green in order to increase her reputation level. If the reason for instances where the bad agent plays green was the other-regarding preferences, then we would observe this behavior in any period - more precisely, including the last period. However, there is no single observation where the bad agent plays green in the last period. Hence, I conclude that reputation concerns triumph over other-regarding preferences in our setup.

Fourth, gradualism has been experimentally shown to be valuable in different environments. Schelling [1960] suggested that dividing contributions into small pieces is better to maximize contributions in public goods games. In this way, counterparts will be able to build trust through time. Duffy et al. [2007] experimentally tested this prediction and showed that contributions are higher in dynamic games than in static games. Weber [2006] showed that by gradually increasing the size of

¹The most prominent earlier works include: Camerer and Weigelt [1988], Neral and Ochs [1992], Anderhub et al. [2002], and Brandts and Figueras [2003]

a group, we could enhance the cooperation within the group. Kamijo et al. [2016] and Ye et al. [2020] showed that gradualism leads to maximum coordination in the lab environment. Kurzban et al. [2008] and Kartal et al. [2021] analyzed repeated trust games. They showed that it is optimal for the sender to implement a gradualist strategy in certain environments, i.e., starting with small trust levels. The participants prefer to start with small trust and gradually increase stakes.

Andreoni and Samuelson [2006] experimentally analyzed a twice-repeated prisoner's dilemma game. They varied the relative stakes of the periods and found that the optimum allocation of stakes is approximately one-third to two-thirds. Starting small in this fashion contributes to cooperation in both periods and increases the players' payoffs. Andreoni et al. [2019] extended that analysis by endogenizing the relative stakes' decision. In their experiments, players chose the relative stakes, and they found that starting small increases the cooperation and social payoff. Moreover, they showed that when one of the players chooses to start small, more cooperation is observed than the exogenous starting small. Hence, the players benefit from starting small, and they allocate stakes strategically to maximize payoffs. Once a player leaves higher stakes to the second period, the threat of ending relation after the first period becomes considerable. The presence of a cooperative type can be manipulated to enhance cooperation. They showed that the subjects could structure the interaction accordingly, even for a short period to build a reputation. I analyze a different stage game than Andreoni et al. [2019]. Moreover, there are other differences in our approach. They provided a prisoner's dilemma and induced subjects' types from the data. In comparison, I define two different types depending on the existence of the bias. I impose preferences on the subjects by setting two different payoff schemes for different types. They focused on the joint payoffs and showed that it is maximized when the interaction starts small. I focus on principal-optimal equilibrium because the principal is the only player with incomplete information.

3.3 The model

For a project of three periods $N = \{1, 2, 3\}$, a principal considers making single-period renewable contracts with an agent. At the beginning of each period $i \in N$, the principal hires the agent with probability $\lambda_i \in [0, 1]$. If $\lambda_i = 0$, the agent is not hired and the project is terminated. If $\lambda_i = 1$, the agent observes the realized state $\theta_i \in \{0, 1\}$ which is not observable by the principal at the time, and takes action $a_i \in \{0, 1\}$. (Each state is equally likely, i.e., $\text{prob}[\theta_i = 0] = \text{prob}[\theta_i = 1] = 0.50$, and independent across the periods.)

The agent can be *good* or *bad* type. The good agent shares the same preferences with the principal, while the bad agent favors $a_i = 1$ regardless of the realized state. Having a prior probability $\rho_1 \in (0, 1)$ that the agent is good type, the principal updates his belief about the agent's type as she observes a_i and θ_i . Then the game proceeds to the next period.

The timing of the stage game in period i can be summarized as follows:

1. The principal decides whether to hire the agent. If $\lambda_i = 0$, then the project is terminated and the payoffs are realized.
2. If $\lambda_i = 1$, nature chooses the state $\theta_i \in \{0, 1\}$.
3. The agent observes θ_i and takes the action $a_i \in \{0, 1\}$

The players maximize their expected payoffs. The expected continuation payoff of a player in period $i \in N$ is given by $V_i = \sum_{j=1}^3 \gamma_j U_j$ where U_i denotes the stage game payoff in period i . $\gamma_i \in \Gamma$ is the stake (weight) of period i where $\Gamma = \{\gamma_1, \gamma_2, \gamma_3\}$ such that $\sum_{i=1}^3 \gamma_i = 1$. Γ is common knowledge and constant for all players.

For $i \in N$, the history is given by $h_i = (\lambda_i, a_i, \theta_i)$ while $H_i = (h_1, \dots, h_{i-1})$ denotes the set of all histories before the principal takes a decision. The principal's strategy in i is given by $\Lambda_i : H_i \rightarrow [0, 1]$ where $\Lambda_i(h)$ is the principal's choice of λ_i . The belief of the principal in period i about the agent's type is defined by the probability of the agent being good, i.e., $\rho_i : H_i \rightarrow [0, 1]$. $I_i = (H_i, \lambda_i, \theta_i)$ is the information

set of the agent in period i . The good agent's (mixed) strategy is given by $\mu_i^{\theta_i} = \text{prob}(a_i = 0|\theta_i)$ and the bad agent's strategy is given by $\nu_i^{\theta_i} = \text{prob}(a_i = 0|\theta_i)$. $\mu_i \& \nu_i : \mathfrak{I}_i \rightarrow \Delta[0, 1]$ where $\mathfrak{I}_i$ is the set of all period i information sets.

For $i \in N$, the stage game payoffs are given in Table 3.1.

Table 3.1: Payoff table.

Good agent (and principal)			Bad agent		
	$\theta_i = 0$	$\theta_i = 1$		$\theta_i = 0$	$\theta_i = 1$
$a = 0$	k	$k - 1$	$a = 0$	$k - 1$	$k - 1$
$a = 1$	$k - 1$	k	$a = 1$	k	k

where k is the stage payoff from the project. Note that $k \in (0, 1/2)$.²

Table 3.2 summarizes the stage game strategy of the agents in period i .

Table 3.2: Strategies table.

	$\theta_i = 0$	$\theta_i = 1$
good agent	0 with probability μ_i^0 1 with probability $1 - \mu_i^0$	0 with probability μ_i^1 1 with probability $1 - \mu_i^1$
bad agent	0 with probability ν_i^0 1 with probability $1 - \nu_i^0$	0 with probability ν_i^1 1 with probability $1 - \nu_i^1$

²Note that if $k \geq 1/2$, the principal unexceptionally hires the agent. For instance, if $\rho_i = 0$ (i.e., the principal knows that the agent is indeed a bad type), the expected payoff of the principal is equal to $k - 1/2 \geq 0$. That is, $\lambda_i = 1$. On the other hand, if $k < 0$, then the principal's expected payoff is negative even if the agent is a good type. I ignore $k = 0$ for convenience since this would imply the principal's hiring the agent only when she is a good type.

In this repeated game with asymmetric information, there are multiple equilibria. We focus on the perfect Bayesian equilibria with Markovian property. In that regard, history matters only in terms of its effect on the agent's reputation. Thus, if the agent has the same reputation record after two different histories, the principal's action is the same in both cases. The principal's decision depends on the past observations only through their effect on the agent's reputation. More formally, for any $i \in N$ and $h, h' \in H_i$, $\rho_i(h) = \rho_i(h')$ implies $\lambda_i(\theta_i, a_i|h) = \lambda_i(\theta_i, a_i|h')$ and $\mu_i(h) = \mu_i(h')$ and $\nu_i^{\theta_i}(h) = \nu_i^{\theta_i}(h')$. I select the principal-optimal equilibrium as the equilibrium candidate whenever there are multiple perfect Bayesian equilibria. Principal-optimal equilibrium corresponds to the equilibrium which yields the highest payoff to the principal among other equilibria.

3.4 Experimental Design

In our experiment, a principal considers making single-period renewable contracts with an agent for a project of three periods $N = \{1, 2, 3\}$. The probability of the agent being a good type and the payoff from the project is given exogenously, i.e., $\rho_1 = 0.1$, and $k = 1/3$, respectively.

There are four treatments based on the varying weights distributions exogenously assigned across the interaction stages between the principal and the agent:

- Treatment 1 (T1): $\Gamma_1 = \{\frac{1}{3}, \frac{1}{3}, \frac{1}{3}\}$.
- Treatment 2 (T2): $\Gamma_2 = \{\frac{1}{7}, \frac{2}{7}, \frac{4}{7}\}$.
- Treatment 3 (T3): $\Gamma_3 = \{\frac{1}{13}, \frac{3}{13}, \frac{9}{13}\}$.
- Treatment 4 (T4): $\Gamma_4 = \{\frac{1}{21}, \frac{4}{21}, \frac{16}{21}\}$.

Figure 3.1 visualizes the treatment allocation of the treatments.

Treatment 1 is the baseline treatment where every period has an equal stake. In the other treatments, the stakes start small and increase at varying speeds through the periods. In Treatment 2, 3 and 4, the stakes get doubled, tripled and quadrupled,

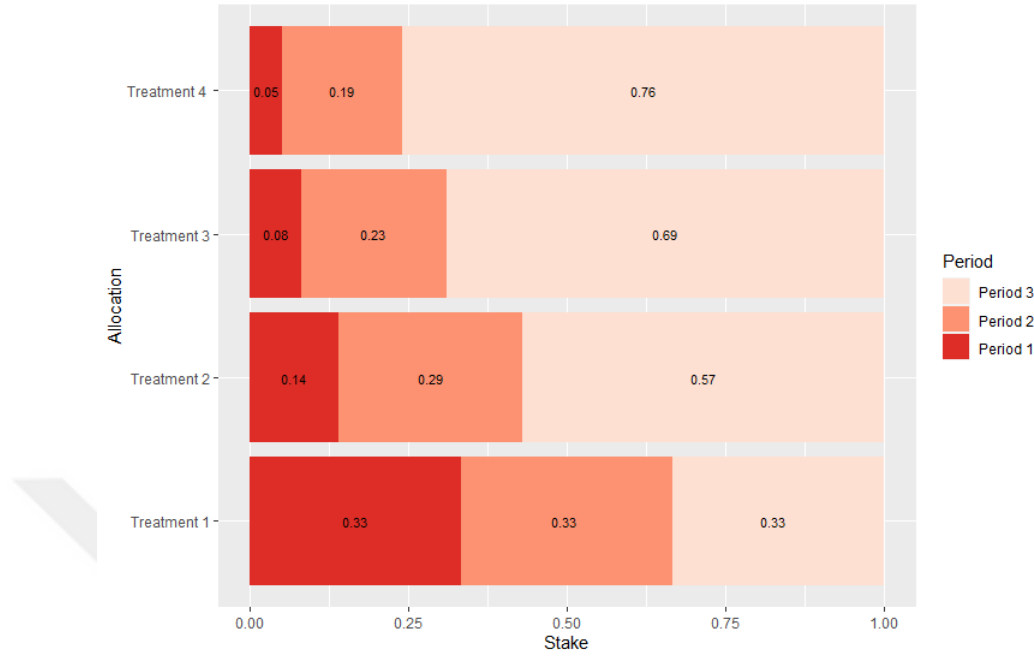


Figure 3.1: The stake allocation for four treatments.

respectively. For each treatment, the payoff tables for each period are given in Appendix C.

It is a between-subject design, i.e., each subject is allocated to only one of the treatments. For each treatment, we run one session with 20 subjects. At the beginning of each session, half of the subjects are randomly selected for the role of the principal. Due to the commonly known $\rho_1 = 0.1$, 1 of the remaining ten subjects is randomly assigned as a good type agent while the rest are assigned as a bad type agent. In the experiment, neutral labels are used for good and bad agents, i.e., Type 1 and Type 2, respectively. The assigned role of each subject remains intact throughout a session.

There are 15 super-games in a session. For each super-game, a principal is matched with an agent (through the stranger matching rule) whose type is not observable by the principal. Each super-game contains maximum 3 periods. The payoff table for each period $i \in N$ of a super-game is common knowledge and remains intact throughout a session. 2 of 15 super-games are randomly selected,

and the subject's earnings from these two periods are paid. For each period in a super-game, the following steps are implemented:

1. The principal is asked about her belief concerning the probability of her matched agent being of Type 1.
2. The principal chooses $\lambda_1 \in \{0, 1\}$, i.e., decides whether to hire the agent for one period.
3. If $\lambda_1 = 0$, the super-game is terminated. If $\lambda_1 = 1$, then $\theta_i \in \{0, 1\}$ is set by the computer. (In the experiment, "0" and "1" are tagged as "green" and "blue", respectively, i.e., $\theta_i \in \{green, blue\}$.)
4. The agent observes the computer's choice and takes an action $a_i \in \{green, blue\}$.
5. The period ends as the principal observes both the agent's action a_i and the computer's choice of θ_i .
6. Unless it is the third period, the super-game proceeds to the next period and the steps 1-5 are repeated.

Prior to her decision in each period, the principal is reminded about the agent's past actions, realized states, and her payoffs. The principal cannot observe the agents' records in the past super-games. That is, reputation building by the agent is defined within each super-game, not throughout super games.

The experiment consists of four modules. In the first module, the subjects are given the instructions followed by an incentivized quiz on their comprehension. The second module constitutes the main body of the experiment with the 15 super-games. In the third module, the risk attitudes of the subjects are elicited via an incentivized Holt-Laury task [Holt and Laury, 2002]. The experiment ends with a survey module in which the subjects are asked to answer questions on demographics in return for a fixed payment.

Given the payoff tables in Appendix C, reflecting the implementation of different stake vectors across treatments, approximating the range of potential earnings across the treatments stands as a design challenge in our experiment. I handle this challenge by employing different exchange rates for each treatment. As the first step in this direction, we constrained each super-game's monetary (Turkish Lira (TL)) payoff to range between -20 TL and 10 TL in all treatments. Given the stake vectors of the treatments, the payoff in each period needs to be a minimum of 273 ECU (Experimental Currency Unit) (as the least common multiple of 3, 7, 13, and 21.) In my aim to have exchange rates such that it is easy for the subjects to calculate the monetary value of the ECU payoff, we employ $1/3$, $1/7$, $1/13$, and $1/21$ as the rates in Treatment 1, 2, 3, and 4, respectively.

Two of the super-games in Module 2 are selected for payment at the end of the experiment. That is, the minimum earning is potentially -40 TL in Module 2. In order to compensate for this potential loss, subjects are given fixed payments and chances to earn money in the other modules of the experiment. As fixed payments, we give a show-up fee of 20 TL and, in Module 4, a survey participation fee of 10 TL. The subjects are not informed about the latter until the beginning of Module 4 in order to avoid an endowment effect. On the other hand, for each correct answer they gave to the quiz questions in Module 1, the subjects received payoffs in ECU worth of 2 TL, i.e., 6, 14, 26, and 42 ECU, respectively, across the four treatments. In Module 3, the subjects could get payoffs at the Holt-Laury task. (Again, the potential payoffs are adjusted to have approximately equal TL values across the treatments.) Table C.5 in Appendix C gives a summary of the payoffs.

The experiments were conducted via oTree [Chen et al., 2016] at the virtual lab of BELIS (Bilgi Economics Lab of Istanbul). The subjects are undergraduate students at various Istanbul Bilgi University departments recruited via ORSEE.

3.5 Theory and Predictions

3.5.1 Reputation Building Behavior

The model predicts that, given any value of θ_i , playing $a_i = 0$ with a positive probability augments the agent's reputation. In other words, $a_i = 1$ never improves the agent's reputation. Hence, the reputation concern corresponds to the agent's incentive to play $a_i = 0$. Proposition 3.1 formalizes the argument.

Proposition 3.1. *In the principal-optimal perfect Bayesian equilibrium,*

1. *There is a reputation incentive for the agent to play $a_i = 0$ when $\theta_i = 0$, i.e.,*

$$\rho_{i+1}(\theta_i = 0, a_i = 0) > \rho_{i+1}(\theta_i = 0, a_i = 1), \quad \forall i \in N \text{ and } \forall h \in H_i.$$

2. $\rho_i \geq \rho_{i+1}(1, 1)$.

Proof. in Appendix E. □

Note that $\rho_{i+1}(0, 0) > \rho_{i+1}(0, 1)$ implies $\mu_i^0 = 1$. The good agent has no incentive to play $a_i = 1$ when she observes $\theta_i = 0$. That is, playing $a_i = 1$ with a positive probability when $\theta_i = 0$ ($\mu_i^0 < 1$) cannot be an equilibrium strategy for the good agent since playing $a_i = 1$ hurts both her stage game payoff and reputation. Hence, $\mu_i^0 = 1$ in the principal-optimal equilibrium. To simplify the notation, I will use μ_i hereafter to refer to μ_i^1 .

Given the strategies defined above, the principal's equilibrium-path belief about the type of the agent in period $i + 1$ ($\rho_{i+1}(\theta_i, a_i)$) - is updated by Bayes' rule as follows:³

$$\rho_{i+1}(0, 0) = \frac{\rho_i}{\rho_i + (1 - \rho_i)\nu_i^0}$$

$$\rho_{i+1}(0, 1) = 0$$

³ $\mu_i = 0$ and $\nu_i^1 = 0$ implies $\rho_{i+1}(1, 1) = \rho_i$.

$$\rho_{i+1}(1, 0) = \frac{\rho_i \mu_i}{\rho_i \mu_i + (1 - \rho_i) \nu_i^1}$$

$$\rho_{i+1}(1, 1) = \frac{\rho_i (1 - \mu_i)}{\rho_i (1 - \mu_i) + (1 - \rho_i) (1 - \nu_i^1)}$$

Prediction A.

1. The principal's belief that the agent is the good type will be the highest after $\{0, 0\}$. More precisely, $\rho_{i+1}(0, 0) \geq \rho_{i+1}(1, 1) > \rho_{i+1}(0, 1)$ ⁴.
2. The principal will hire the agent most frequently after $\{0, 0\}$ followed by $\{1, 1\}$ and $\{0, 1\}$: $\lambda_{i+1}(\theta_i = 0, a_i = 0) > \lambda_{i+1}(\theta_i = 1, a_i = 1) > \lambda_{i+1}(\theta_i = 0, a_i = 1)$.

I am looking for the principal-optimal allocation of γ_i when there is a sufficiently small probability that the agent is the good type. My theoretical model predicts the principal-optimal $\gamma_i = \frac{1-k}{1-k^n} k^{n-i}$. I test $k = \frac{1}{3}$ in our experiment; hence, the predicted optimal stake allocation is as follows: $\gamma_i = \frac{3^{i-1}}{27}$. This corresponds to $\Gamma_3 = \{\frac{1}{13}, \frac{3}{13}, \frac{9}{13}\}$. For this allocation of stakes, the bad agent is indifferent between $a_i = 0$ and $a_i = 1$ in the first and second periods, while the good agent prefers $a_i = 0$. The principal hires the agent after observing $a_i = 0$ and does not hire after $a_i = 1$. In this setting, the agent's reputation increases gradually through her interaction with the principal. I assume that the principal sets the relative stakes to implement the principal-optimal allocation. This assumption does not alter the equilibrium analysis as any third party can set the stakes whose objective is to maximize the principal's expected continuation payoff.

The bad agent plays $a_i = 0$ with a positive probability in the first and the second periods when $\theta_i = 0$. The mixing probability of the agent must stay within a

⁴In the equilibrium we experimentally analyze, the observation $\{a_i = 0, \theta_i = 1\}$ is never realized. The principal's belief $\rho_{i+1}(1, 0)$ is free. Thus, I do not include this observation in the predictions. Please refer to Remark 1 in Appendix for a detailed argument.

certain range to meet the equilibrium beliefs. The reputation incentives of the bad agent have a two-dimensional effect on the expected payoff of the principal. The principal's period i payoff is increasing in ν_i^0 . This corresponds to the "discipline effect" of reputation incentives as defined by Morris [2001]. In addition, for large enough k , the expected continuation payoff of the principal is also increasing in ν_i^0 . I focus on large k values in this analysis. We use principal-optimality as the equilibrium selection criteria and for determining ν_i^0 . As the expected continuation payoff of the principal is increasing in ν_i^0 , the agent will play $a_i = 0$ with the highest probability satisfying the equilibrium conditions.

The following proposition formalizes the principal-optimal equilibrium.

Proposition 3.2. *If we start with a sufficiently low initial reputation ($\rho_1 \ll (1 - 2k)^2$) and $k \geq \frac{1}{4}$, $\gamma_i = \frac{1-k}{1-k^3}k^{3-i}$ is the principal-optimal allocation of stakes and the following strategies constitute the principal-optimal perfect Bayesian equilibrium:*

$$\nu_1^0 = \frac{4k\rho_1(1-k)}{(1-\rho_1)(1-2k)^2}, \nu_2^0 = \frac{2k\rho_1}{(1-\rho_1)(1-2k)}$$

$$\nu_1^1 \& \nu_2^1 = 0$$

$$\mu_1^1 \& \mu_2^1 = 0$$

$$\lambda_{i+1}(0,0) = 1, \lambda_{i+1}(\theta_i,1) \& \lambda_{i+1}(0,1) = 0 \text{ for } \theta_i \in \{0,1\} \text{ and } i \in \{2,3\}.$$

Proof. Appendix E. □

The principal's ex-ante expected payoff is the highest if the stakes are allocated according to $\Gamma_3 = \{\frac{1}{13}, \frac{3}{13}, \frac{9}{13}\}$. The principal continues to hire the agent as long as $\{\theta_i, a_i\} = \{0, 0\}$ is observed in every period.

Treatment 1: Equal stakes ($\Gamma_1 = \{\frac{1}{3}, \frac{1}{3}, \frac{1}{3}\}$)

The principal's expected continuation payoff is negative in this allocation, more precisely, -0.028 . Hence, the risk-neutral principal is expected never to hire the

agent. However, contrary to this theoretical prediction, hiring behavior is likely to be observed in experiments due to home-made priors and/or risk preferences ⁵. Thus, for the equilibrium, I assume $\lambda_1 = 1$.

Equal stake allocation implies no reputational incentive to exist in any period. ⁶ If $\{\theta_1, a_1\} = \{0, 0\}$ is observed in the first period, then the principal knows that the agent is the good type. The principal hires the agent in the remaining periods irrespective of the observation. Otherwise, the principal never hires the agent.

The following proposition formalizes the unique perfect Bayesian equilibrium.

Proposition 3.3. *Set $\rho_1 = 0.1$ and $k = \frac{1}{3}$ where the allocation of stakes is given by $\Gamma_1 = \{\frac{1}{3}, \frac{1}{3}, \frac{1}{3}\}$. Assuming $\lambda_1 = 1$, the following assessment constitutes the unique perfect Bayesian equilibrium:*

$$\nu_i^{\theta_i} = 0 \text{ for } \theta_i \in \{0, 1\} \text{ and } \mu_i = 0 \forall i \in \{1, 2, 3\}$$

$$\lambda_2(0, 0) = 1, \lambda_2(1, 0) = \lambda_2(1, 1) = \lambda_2(0, 1) = 0$$

Proof. Appendix E. □

Treatment 2: Stakes start small and increase gradually ($\Gamma_2 = \{\frac{1}{7}, \frac{2}{7}, \frac{4}{7}\}$)

The principal has a positive expected continuation payoff: 0.0024. The agent has a reputation incentive only in the first period. Compared to Treatment 1, there are stronger reputational incentives. That is, observing $\{\theta_1, a_1\} = \{0, 0\}$ leads the agent's reputation to evolve slower in Treatment 2 than in Treatment 1. ⁷

Proposition 3.4. *Set $\rho_1 = 0.1$ and $k = \frac{1}{3}$ where the allocation of stakes is given by $\Gamma_2 = \{\frac{1}{7}, \frac{2}{7}, \frac{4}{7}\}$. The following assessment constitutes the principal-optimal perfect*

⁵See Camerer and Weigelt [1988] for the discussion on the home-made priors.

⁶The evolution of reputation is following in the above equilibrium: $\rho_2(0, 0) = 1$, $\rho_2(1, 1) = \rho_1$ and $\rho_2(1, 0) = \rho_2(0, 1) = 0$

⁷The evolution of reputation is as follows in the equilibrium: $\rho_2(0, 0) = \frac{1}{3}$, $\rho_2(1, 1) = \rho_1$, $\rho_2(1, 0) = \rho_2(0, 1) = 0$, $\rho_3(0, 0) = 1$, $\rho_3(1, 1) = \rho_2$, $\rho_3(1, 0) = \rho_3(0, 1) = 0$.

Bayesian equilibrium:

$$\nu_1^0 = \frac{2\rho_1}{1-\rho_1} = \frac{2}{9}, \nu_2^0 = 0$$

$$\nu_i^1 = \mu_i = 0 \forall i \in \{1, 2, 3\}.$$

$$\lambda_2(0, 0) = 1, \lambda_2(1, 0) = \lambda_2(1, 1) = \lambda_2(0, 1) = 0$$

$$\lambda_3(0, 0) = 1, \lambda_3(1, 1) = \frac{1}{2}, \lambda_3(1, 0) = \lambda_3(0, 1) = 0$$

Proof. Appendix E. □

Treatment 3: Stakes start small and increase at medium speed ($\Gamma_3 = \{\frac{1}{13}, \frac{3}{13}, \frac{9}{13}\}$)

The principal has an expected continuation payoff of 0.0218. The agent has a reputation incentive in all periods except the last one. Compared to Treatment 1 and Treatment 2, the reputation will evolve slower. Due to Proposition 3.2, Γ_3 is the optimal allocation of the stakes for the principal. As there are more reputation incentives, the reputation of the agent will evolve slower in this setup.⁸

Proposition 3.5. *Set $\rho_1 = 0.1$ and $k = \frac{1}{3}$ where the allocation of stakes is given by $\Gamma_3 = \{\frac{1}{13}, \frac{3}{13}, \frac{9}{13}\}$. The following assessment constitutes the principal-optimal equilibrium:*

$$\nu_1^0 = \frac{8\rho_1}{1-\rho_1} = \frac{8}{9}, \nu_2^0 = \frac{2\rho_2}{1-\rho_2}$$

$$\nu_i^1 \& \mu_i = 0 \text{ for } i \in \{1, 2, 3\}$$

$$\lambda_i(0, 0) = 1, \lambda_i(1, 0) = \lambda_i(1, 1) = \lambda_i(0, 1) = 0 \text{ for } i \in \{2, 3\}$$

Proof. Appendix E. □

⁸The evolution of reputation in the equilibrium: $\rho_2(0, 0) = \frac{1}{9}, \rho_2(1, 0) < \frac{1}{9}, \rho_2(0, 1) = 0, \rho_2(1, 1) = \rho_1 = \frac{1}{10}, \rho_3(0, 0) = \frac{1}{3}, \rho_3(1, 0) < \frac{1}{3}, \rho_3(0, 1) = 0, \rho_3(1, 1) = \rho_2$.

Treatment 4: Stakes start small and increase at high speed ($\Gamma_4 = \{\frac{1}{21}, \frac{4}{21}, \frac{16}{21}\}$)

This allocation of stakes clearly implies high reputation incentives. That is, if the principal hires the agent after observing $(0, \theta_i)$, then the agent's best response will be to play $a_i = 0$ with probability 1. Then, the reputation is not updated. As the initial reputation is sufficiently low, this behavior can not constitute an equilibrium. Thus, the principal plays a mixed strategy. The principal gets the expected payoff of 0.013 in this setup.⁹

Proposition 3.6. *Set $\rho_1 = 0.1$ and $k = \frac{1}{3}$ where the allocation of stakes is given by $\Gamma_4 = \{\frac{1}{21}, \frac{4}{21}, \frac{16}{21}\}$. The following assessment constitutes the principal-optimal equilibrium:*

$$\nu_1^0 = \frac{8\rho_1}{1-\rho_1} = \frac{8}{9}, \quad \nu_2^0 = \frac{2\rho_2}{1-\rho_2}$$

$$\nu_i^1 \& \mu_i = 0 \text{ for } i \in \{1, 2, 3\}$$

$$\lambda_2(0, 0) = \lambda_3(0, 0) = \frac{3}{4}$$

$$\lambda_i(0, 1) = \lambda_i(1, 0) = \lambda_i(1, 1) = 0 \text{ for } i \in \{2, 3\}$$

Proof. Appendix E. □

Prediction B. *The smaller the interaction starts, the (weakly) higher are the reputational incentives.*¹⁰

⁹The evolution of reputation is as follows in the above equilibrium: $\rho_2(0, 0) = \frac{1}{9}$, $\rho_2(0, 1) < \frac{1}{9}$, $\rho_2(1, 0) = 0$, $\rho_2(1, 1) = \rho_1$, $\rho_3(0, 0) = \frac{1}{3}$, $\rho_3(0, 1) < \frac{1}{3}$, $\rho_3(1, 0) = 0$, $\rho_3(1, 1) = \rho_2$.

¹⁰In Treatment 3, we choose the highest ν_i^0 (that is supported in the equilibrium) to define the principal-optimal equilibrium. Highest ν_i^0 corresponds to the lowest $\rho_{i+1}(0, 0)$. Hence, we choose the highest probability with which the agent invests in reputation and the slowest evolution of stakes. In Treatment 4, on the other hand, the model predicts a unique strategy of the agent. The reason is that the mixing strategy of the agent needs to be so that the principal is indifferent between hiring or not, and hence, the principal mixes his own strategies. As I show that the agent's strategy is the same in Proposition 3.5 and 3.6, I can conclude that the reputation incentives will be weakly higher in Treatment 4.

3.5.2 Payoffs

The theoretical model predicts that if the agent's initial reputation is sufficiently low, the principal-optimal allocation of relative stakes is such that the bad agent is indifferent between $a_i = 0$ and $a_i = 1$ in each period. The stakes increase gradually in the equilibrium; hence, our findings comply with the literature on gradualism. If the interaction starts with a higher stake and the stakes evolve slower (Treatment 2), then the agent will not have sufficient reputation incentives. On the other hand, if the interaction starts very small and the stakes increase rapidly (Treatment 4), the reputation incentive will be so high that the principal cannot update his belief regarding the agent's type. Hence, gradualism with the "right" speed is expected to lead to the highest payoff.

Prediction C. *In terms of the expected payoff of the principal, the stake allocations are ranked as follows: $\Gamma_3 > \Gamma_4 > \Gamma_2 > \Gamma_1$.*¹¹

According to this prediction, starting the interaction small is better for the principal if the initial reputation is sufficiently low. However, if the stakes do not increase at the right speed, there will be an excess or a deficit of reputation incentive. Hence, the result will be sub-optimal for the principal. If the interaction starts smaller than the equilibrium stake and evolves faster, there will be high reputational incentives. However, three types of gradualism are preferred to starting small.

3.6 Results

I ran the experimental sessions in December 2020. For each treatment, we run one session with 20 subjects. The average duration of a session was approximately 100 minutes. The subjects received 52 TL on average as experimental earnings.

In the following subsections, I present our experimental findings on the principal and agent behavior as well as their incentives and realized payoffs.

The subject's behavior is in accordance with the predictions of theory to some extent. The reputation incentives and the principal's average payoff are observed

¹¹ $U^P(T3) = 0.0218$, $U^P(T4) = 0.013$, $U^P(T2) = 0.0024$, $U^P(T1) = -0.028$

to vary across treatments. As expected, the smaller the interaction starts, the more reputation incentives are there.

The first thing to note is that the good agent never plays the action different from the state. They follow the babbling strategy.

The main divergence from the predictions occurs when the bad agent plays action 1 when the state is equal to 0. The good agent has no motivation to act accordingly; hence the principal is expected to not hire the agent after observing such an outcome. However, the principal hires the agent in 64.6% of all such occurrences. This behavior can be explained by risk incentives and subjective beliefs about the next period's state in the following subsections.

3.6.1 Agent Behavior and Reputation Incentives

I measure the agent's reputation incentive by the frequency of the bad agent playing $a_i = 0$, i.e., ν_i^0 . It should be noted that, theoretically, there remains no reputation incentive in the last (third) period, so we expect $\nu_3^{\theta_3} = 0$ in all treatments. Our experimental data confirm this prediction: none of the bad agents plays $a_3 = 0$. Hence, we can deduce that $\nu_i^0 > 0$ for $i = \{1, 2\}$ stems from reputation incentives rather than some other motives of the bad agents (such as altruistic preferences, for instance.) For the analysis, I focus on the first-period reputation incentives since the second-period reputation incentives depend not only on the allocation of stakes but also on the first-period outcomes.¹²

My theoretical model predicts two possible equilibrium behavior when $\theta_i = 1$: The first is the babbling equilibrium behavior at which both types of agents play $a_i = 1$. The principal's out-of-equilibrium path belief $\rho_{i+1}(1, 0)$ must be sufficiently small so that the agent has no reputation incentive to play $a_i = 0$. I observe babbling equilibrium behavior in almost all treatments in our experimental data. The second equilibrium behavior is the reputation-building behavior which I seldom observe only for the bad agent in Treatment 3. In the 16 of total 72 occurrences of $\theta_1 = 1$, i.e., $\nu_2^1 = 0.22$, and $\rho_2(1, 0)$ is reported to be 49.1 on average. Given these results, I

¹²Our results indicate a similar yet insignificant trend for the second-period incentives.

focus on the case $\theta_1 = 0$ to analyze reputation incentives.

In Table 3.3, I report the observed ν_1^0 and the outcome of logit regression. Our experimental results confirm the Prediction B that the smaller the initial stake is, the stronger the reputation incentives are, i.e., $\nu_i(T1) < \nu_i(T2) < \nu_i(T3) < \nu_i(T4)$ for $i = \{1, 2\}$.¹³

3.6.2 Beliefs

When a period starts, the principal is asked about his belief in the agent's probability of being the good type. At the beginning of the first period, they are informed that there is a 10% chance of matching with a good agent. 85% of the principal reported a belief of 10% in the first period, and the mean belief is 13%, which is reasonably close to the objective probability of the agent being the good type. Overall, principals tend to report higher beliefs (compared to Bayesian updating). One such instance is as follows. After observing action blue and state green, we expect the principal's belief to be zero. However, the average observed belief in period 2 after such observation is 10%. This phenomenon can be explained by the homemade prior beliefs of the principals about the agent's type.

If $a_i = 1$ and $\theta_i = 0$ is observed in any period, then it implies that the agent is a bad type, i.e., $\rho_{i+1}(0, 1) = 0$. If $a_1 = 1$ and $\theta_1 = 0$ is observed, 87% of the principals reported belief equal or below 10% (54,5% being equal to zero). Even though it is significantly away from the expected belief of 0, we observe that $\{\theta_1, a_1\} = \{0, 1\}$ decreases the agent's reputation. If we define $\Delta\rho_{21}$ as $\rho_2 - \rho_1$, the mean $\Delta\rho_{21}$ is equal to $-3, 2$, and 91, 84% of observed $\Delta\rho_{21}$ values are non-positive. In percentage terms, the average change in beliefs is -112% . $\Delta\rho_{32}(0, 1)$ has also a negative mean of -4.5 .

As expected, $\{\theta_1, a_1\} = \{0, 0\}$ increases the agent's reputation for all treatments.

¹³In Treatment 1, our model predicts a rational bad agent to have no incentive to differentiate between her actions depending on the state. However, if the agent has other-regarding preferences, then we expect to observe the bad agent playing $a_i = 0$ more frequently when $\theta = 0$ compared to $\theta = 1$.

$\Delta\rho_{21}(0,0)$ has an average of 28.65 across all treatments where the average $\rho_{21}(0,0) = 42.76$. We expect it to be higher in the treatments with lower reputational incentives; however, no such trend is observed. $\Delta\rho_{32}(0,0)$ has a mean of 18, which again holds with our expectation of reputation concerns at force.

The principals seem to be more optimistic than the theory predicts when $\{\theta_1, a_1\} = \{1, 1\}$ is observed. At best, we expect no reputation loss, i.e., no reputation update. However, the average $\Delta\rho_{21}(1,1)$ is equal to 5.3 across all treatments, significantly different from 0. On the other hand, $\Delta\rho_{32}(1,1)$ is equal to 4.9, which is not different from 0 at 99% significance level. Table 3.4 and Figure 3.2 summarize the change in belief results.¹⁴ The belief updating is far away from Bayesian updating. However, the sign of change in beliefs follows what theory predicts. Table 3.5 and Graph 3.3 show that Prediction A is confirmed.¹⁵

In addition to the principal's belief about the agent's type, there seems to be another subjective belief in force. In the survey, at the end of the experiment, the subjects were asked about their beliefs in the next period's state being 0, given that the current period's state is 0. Most of them reported beliefs ($\pi(\theta_{i+1} = 0 | \theta_i = 0)$) less than 50%. This belief, alongside risk behavior, can explain why some of the principals continued hiring the agent after observing $\{\theta_1, a_1\} = \{0, 1\}$. Even if they knew that there is a small or no probability that the agent is a good type, they subjectively believed that the probability of two consecutive period's states being 0 is sufficiently low; hence they hired the agent. Furthermore, they reported lower belief on the third period's state being 0 given state 0 is observed in the first two periods ($\pi(\theta_3 = 0 | \theta_1 = 0 \& \theta_2 = 0)$).

¹⁴Although OLS results give us more insight, a closer look suggests that the change in belief data is not normally distributed. As a complementary perspective to look at the data, we run a non-parametric test. We apply Wilcoxon rank-sum (Mann-Whitney) test on belief change and obtain exactly the same outcome. More precisely, $\Delta\rho_{21}(0,0) >^{***} \Delta\rho_{21}(1,1) >^{***} \Delta\rho_{21}(0,1)$.

¹⁵Although OLS results give us more insight, a closer look suggests that the belief data is not normally distributed. As a complementary perspective to look at the data, we run a non-parametric test. We apply the Wilcoxon rank-sum (Mann-Whitney) test on beliefs and obtain the same outcome. More precisely, $\rho_2(0,0) >^{***} \rho_2(1,1) >^{***} \rho_2(0,1)$.

Note that $\pi(\theta_{i+1} = 0 | \theta_i = 0)$ first-order stochastically dominates $\pi(\theta_3 = 0 | \theta_1 = 0 \& \theta_2 = 0)$, which is no surprise (Table 3.6). The subjects form a subjective belief in the state such that there is a less than 1/2 probability of observing state 0 in two consecutive periods. Observing 0 in all three periods is expected to be even less likely.

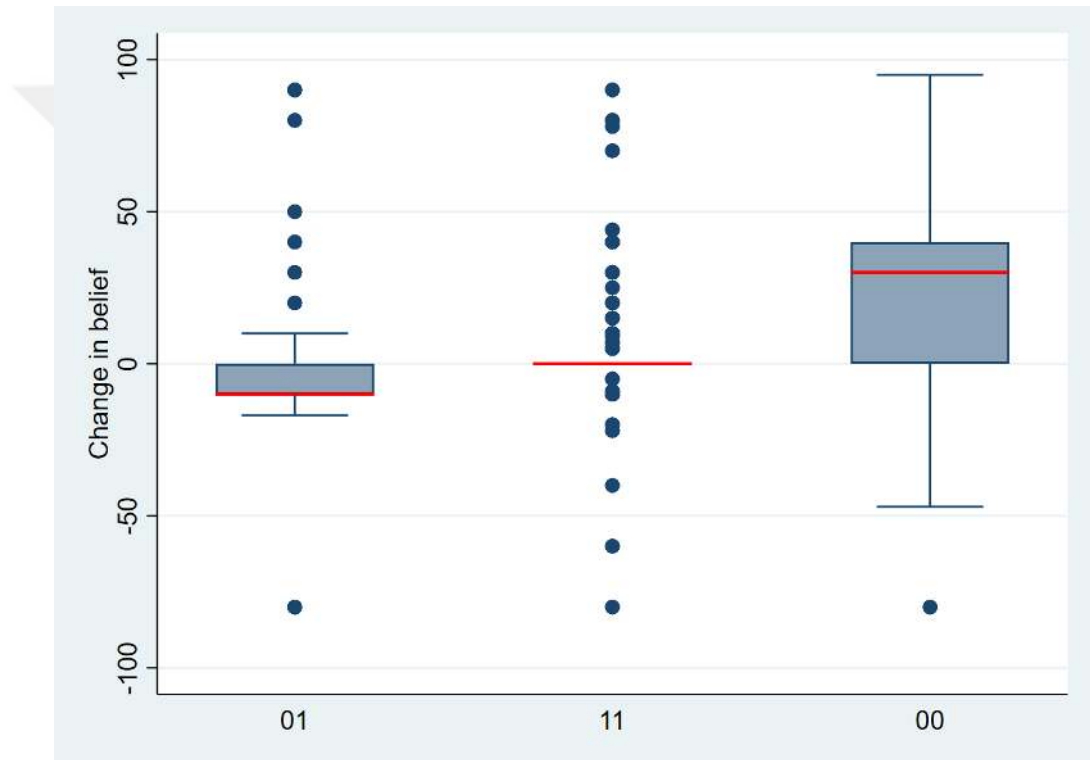


Figure 3.2: Change in beliefs. (01 corresponds to the observation of state 0 and action 1.)

3.6.3 Principal Behavior

The principal hires the agent if he expects a non-negative payoff. A risk-neutral principal with no subjective belief in the state does not hire the agent if he knows

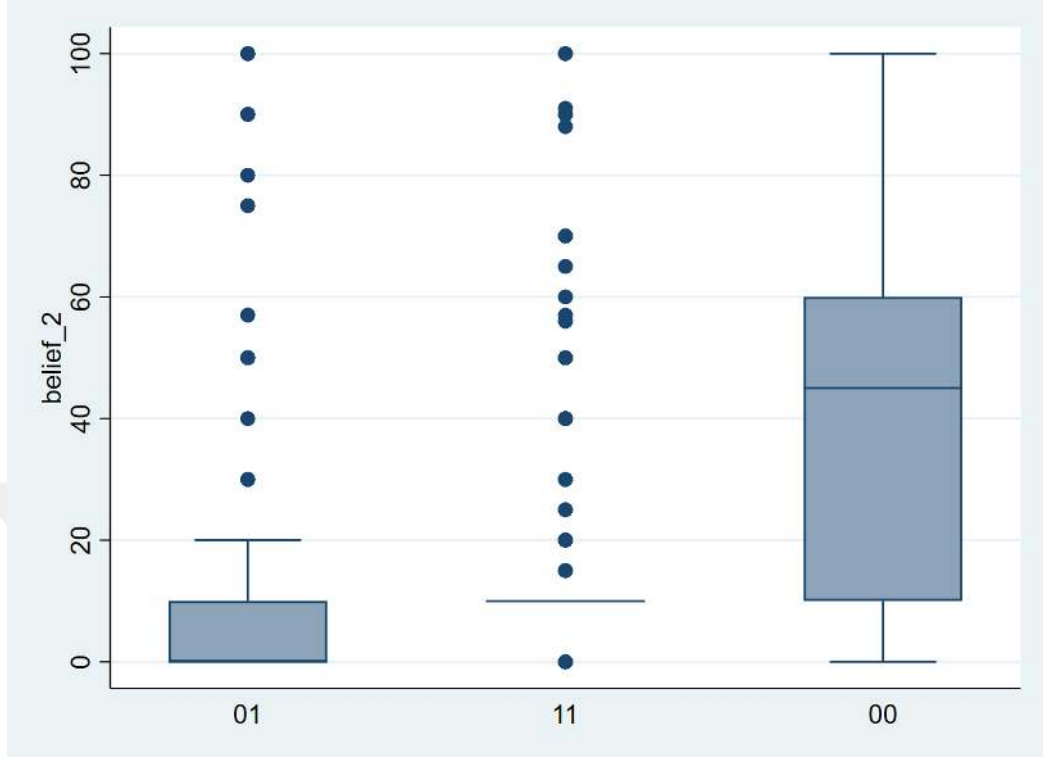


Figure 3.3: Second period beliefs. (01 corresponds to the observation of state 0 and action 1.)

the agent is a bad type with a high probability. In the particular case where the agent plays action 1 while the state is 0, the principal should not hire the agent. However, we observe a positive hiring rate after such an occurrence. In general, the hiring rates are higher than the theory's prediction. Homemade priors about either the type of the agent or the state in the next period combined with risk preferences may lead the principal to hire the agent even in circumstances where a risk-neutral principal is expected not to hire the agent.

The behavior pattern, on the other hand, is in line with the theory's predictions. The higher the principal's reported belief in the agent's type is, the higher the hiring rate is observed. Moreover, the hiring rate is the highest after $\{\theta_1, a_1\} = \{0, 0\}$. The agent is the least hired after $\{\theta_1, a_1\} = \{1, 0\}$. As discussed earlier, we expected this rate to be 0. Although this prediction is not met, the hiring rate is still sufficiently low compared to other observations in the first and second periods.

Table 3.7 summarizes the logistic regression results. Prediction A.2 receives full support, as seen from the table.

3.6.4 Payoff Comparison

The principal prefers to start the interaction small if the agent is a bad type with a sufficiently high probability. The intuition is that the principal can benefit from the reputation incentive of the agent. Morris [2001] defines three effects of the reputation incentives. First is the “discipline effect,” which is responsible for the biased agent to play the action matching the state in the initial periods. It is the most prominent effect in our setup. The smaller the interaction starts, the higher the reputation incentives, hence the higher the discipline effect is. The second is the “sorting effect,” which enables the principal to update his belief about the agent’s type. In the babbling equilibrium - which is observed in most of the games - this effect is reversed. The principal can not update his belief after $\theta_i = 1$ if both agents play action 1 with a probability of 1. On the other hand, if $\theta_i = 0$, the principal would know the agent’s type in the absence of reputation incentives. In such an environment, the good agent will play $a_i = 0$, and the bad agent will play $a_i = 1$. The strategies imply $\rho_{i+1}(0,0) = 1$ and $\rho_{i+1}(0,1) = 1$ by Bayes’ rule. The third effect is the “political correctness effect,” which arises when the good agent plays action 0 to increase reputation when the state is 1. However, we observe that the good agent always plays the action matching the true state. Hence, this effect also does not apply to the setting of this chapter.

To compare the principals’ payoff across treatments, I collect the strategies across all periods and run a simulation based on the strategies. Approximately 1000 observations per treatment are obtained. Two simulations are conducted. In the first one, all the strategies are gathered. I observe that the payoffs in Treatment 2 and Treatment 4 are significantly higher than those in Treatment 1. There is no other significant difference. With this data, I can conclude that starting small is preferred over the equal-stakes treatment. (Table 3.8)

In the second simulation, I focus on the strategies of slightly risk-neutral subjects.

I have a self-assessment risk question and an incentivized version of the Holt-Laury task [Holt and Laury, 2002] to measure the risk preferences. In this simulation, I exclude the subjects in the extremes in either of these two risk assessment methods. More precisely, the subjects who reported self-riskiness less than 2 and higher than 8 (on a scale between 0 and 10) and those whose switching point is greater than 9 or less than 2 at the Holt-Laury task are excluded. The reason for making a weak restriction such as relative risk neutrality is not to lose many observations. Here I observe all three starting small treatments to lead to a higher payoff to the principal than the equal stakes treatment. No significant difference among the starting small treatments is observed. As the expected payoffs in these treatments are close, it is difficult to identify a significant difference among them through experiments. I conclude that Prediction C is partially supported.

I run an additional simulation where the “consistency” restriction is added to the above-defined slightly risk-neutral agents. Consistency is defined as not hiring the agent if action 1 is observed given the state is 0. The reason for using the observations rather than beliefs to define consistency is that a considerable amount of principals report a positive belief after such outcomes. It may stem from home-made priors or belief that a good agent may have made a mistake. Moreover, another belief that may lead the principal to hire the agent is the subjective belief in the next period’s state. As discussed earlier, the subjects tend to believe the next period’s state will be 1 with a probability higher than 0.5 given $\theta = 0$ is observed in the current period. Hence, by focusing on the hiring outcome for the restriction, I can control for both these beliefs. I do not exclude the whole super-game in which $\lambda_{i+1}(0, 1) = 1$ for some i . Rather, I include the super-game strategies in the simulation and make correction for $\lambda_{i+1}(0, 1)$ as $\lambda_{i+1}(0, 1) = 0$. In this simulation, starting small treatments lead to higher payoffs again. Moreover, Treatment 3 leads to a significantly higher payoff compared to Treatment 2. As the expected payoff in Treatment 3 and 4 is close, I cannot conclude on the comparison.¹⁶

¹⁶ $T4 > T3 >^{**} T2 >^{***} T1$, $T4 >^{***} T1$, where $^*p < 0.10$, $^{**}p < 0.05$, $^{***}p < 0.01$.

3.7 Conclusion

Repeated play becomes valuable in settings of uncertainty and asymmetric information. The agent holds reputation incentives in a repeated interaction between a principal and an agent. A high reputation level is vital for the agent because it determines the principal's future hiring decisions. I show that reputation incentives depend on the relative importance of periods of interaction. Investing in reputation may be costly. In such circumstances, the agent faces a trade-off between taking the costly action that will increase her reputation and the action that maximizes the period payoff. The future interaction's stake should be sufficiently high for the agent to invest in reputation. The main focus of this chapter is whether the principal benefits from the reputation incentives of the agent. I theoretically and experimentally show that it is optimal for the principal to start the interaction small in an environment with a low probability that the agent is a good type. Reputation incentives motivate the bad agent to play the action matching the real state.

This chapter shows that we can measure reputation incentives via online experiments. I observe that principals and agents comprehend the reputation incentives and act accordingly. The higher the future stakes are, the higher the reputation incentives are. In the baseline treatment, each period receives an equal stake. It corresponds to the classic repeated interaction without discounting. In this setup, the importance of decisions does not vary across the periods. If the interaction lasts for three periods, I expected there to be no reputation incentives. I observed small reputation incentives. In the starting small treatments, higher reputation incentives are observed. More precisely, the faster the stakes increase, the higher the reputation concerns.

Principals' belief updates and hiring frequencies are also observed to align with the theory's predictions. The hiring frequency is the highest after the computer and the agent choose (green, green). As expected, the principal hires the agent the least after the agent chooses blue when the computer's choice is green. This outcome signals that the agent is a bad type for sure. However, we observed a positive hiring

rate. I explain this with two phenomena. First, our theoretical results rely on the risk neutrality assumption. The subjects might want to take a risk and hire the agent even if they know that she is a bad type. Second, the subjects might put a probability less than $1/2$ to the possibility of the computer choosing green in two consecutive periods. If a principal knows that the agent is a bad type and believes that the computer will choose blue with a sufficiently high probability in the next period, it is rational to hire the agent. Answers to the related survey questions justify this statement. Many subjects reported that they believed the computer would choose green with a probability of less than $1/2$ in the next period, given that it has chosen green in the current period.

I conclude that starting the interaction small improves the principal's payoff in an environment with high uncertainty. It helps the principal in three ways. First, as the principal is cautious about the agent, starting with small stakes means lower risk. Second, starting small and gradually increasing the stakes “disciplines” the bad agent to choose the same color as the computer. Lastly, the principal can observe the agent's behavior and decrease the uncertainty in the further periods where there are higher stakes.

Table 3.3: Reputation incentives across treatments.

	(1)	(2)
	ν_1^0	rep
Treatment 1	0.098 (0.04)	
Treatment 2	0.324 (0.06)	0.83*** (0.27)
Treatment 3	0.492 (0.06)	1.27*** (0.27)
Treatment 4	0.81 (0.05)	2.18*** (0.28)
Constant		-1.29*** (0.22)
N		263

(1) Observed average probabilities of the bad agent playing $a_1 = 0$ when $\theta_1 = 0$. Standard errors in parentheses.

(2) Probit regression of probability of the bad agent playing $a_1 = 0$. Standard errors in parentheses. Dependent variable is reputation. It takes value 1 if the bad agent plays $a_1 = 0$, 0 otherwise given that $\theta_1 = 0$.

* $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table 3.4: Belief change.

$\{\theta_1, a_1\}$	$\Delta\rho_{21}(a_1, \theta_1)$
$\{0, 0\}$	28.65*** (2.6)
$\{1, 1\}$	5.31*** (1.1)
$\{0, 1\}$	-3.09*** (1.39)
Constant	No
Subject fixed effects	Yes
R-squared	0.25

OLS regression results of difference between second period belief and the first period belief on the first period's observation. Standard errors in parentheses.

* $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table 3.5: Belief update.

$\{\theta_1, a_1\}$	ρ_2
$\{0, 0\}$	42.77*** (2.33)
$\{1, 1\}$	18.61*** (1.27)
$\{0, 1\}$	10.61*** (1.84)
Constant	No
Subject fixed effects	Yes
R-squared	0.534

OLS regression results of second period reputation on the first period's observation.

Standard errors in parentheses.

* $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table 3.6: Frequency of subjective believes in the state.

	0.5	0.4-0.5	0.3-0.4	0.2-0.3	< 0.2	N
$\pi(\theta_3 = 0 \theta_1 = 0 \& \theta_2 = 0)$	45%	10%	18.75%	10%	16.25%	80
Cumulative	45%	55%	73.75%	83.75%	100%	
$\pi(\theta_{i+1} = 0 \theta_i = 0)$	55%	10%	15%	10%	10%	80
Cumulative	55%	65%	80%	90%	100%	

$\pi(\theta_{i+1} | x)$ denotes reported subjective belief in θ_{i+1} given x is realized.

Table 3.7: Hiring behavior.

$\{\theta_{i-1}, a_{i-1}\}$	Odds Ratio
$\{0, 0\}$	13.7*** (3.17)
$\{1, 1\}$	5.25*** (0.61)
$\{0, 1\}$	1.83*** (0.22)
N	294

Logistic regression. Standard errors in parentheses. Dependent variable is hiring. It takes value 1 if the agent is hired, 0 otherwise.

* $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table 3.8: Payoff comparison.

General	$T4 > T2 > T3 > T1$, $T2 >^{**} T1$, $T4 >^{**} T1$
Slightly risk neutral	$T4 > T3 > T2 >^{***} T1$, $T4 >^{***} T1$, $T3 >^* T1$,

Wilcoxon rank-sum (Mann-Whitney) test.

* $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

BIBLIOGRAPHY

- Anderhub, V., Engelmann, D., and Güth, W. (2002). An experimental study of the repeated trust game with incomplete information. *Journal of Economic Behavior & Organization*, 48(2):197–216.
- Andreoni, J., Kuhn, M. A., and Samuelson, L. (2019). Building rational cooperation on their own: Learning to start small. *Journal of Public Economic Theory*, 21(5):812–825.
- Andreoni, J. and Samuelson, L. (2006). Building rational cooperation. *Journal of Economic Theory*, 127(1):117–154.
- Atakan, A., Koçkesen, L., and Kubilay, E. (2020). Starting small to communicate. *Games and Economic Behavior*, 121:265–296.
- Brandts, J. and Figueras, N. (2003). An exploration of reputation formation in experimental games. *Journal of Economic Behavior & Organization*, 50(1):89–115.
- Camerer, C. and Weigelt, K. (1988). Experimental tests of a sequential equilibrium reputation model. *Econometrica: Journal of the Econometric Society*, pages 1–36.
- Celetani, M., Fudenberg, D., Levine, D. K., and Pesendorfer, W. (1996). Maintaining a reputation against a long-lived opponent. *Econometrica: Journal of the Econometric Society*, pages 691–704.
- Chen, D. L., Schonger, M., and Wickens, C. (2016). otree—an open-source platform for laboratory, online, and field experiments. *Journal of Behavioral and Experimental Finance*, 9:88–97.
- Duffy, J., Ochs, J., and Vesterlund, L. (2007). Giving little by little: Dynamic voluntary contribution games. *Journal of Public Economics*, 91(9):1708–1730.

- Ely, J., Fudenberg, D., and Levine, D. K. (2009). When is reputation bad? In *A Long-Run Collaboration On Long-Run Games*, pages 177–205. World Scientific.
- Ely, J. C. and Välimäki, J. (2003). Bad reputation. *The Quarterly Journal of Economics*, 118(3):785–814.
- Fudenberg, D. and Levine, D. K. (1989). Reputation and equilibrium selection in games with a patient player. *Econometrica: Journal of the Econometric Society*, pages 759–778.
- Fudenberg, D. and Levine, D. K. (1992). Maintaining a reputation when strategies are imperfectly observed. *Review of Economic Studies*, 59(3).
- Ghosh, P. and Ray, D. (1996). Cooperation in community interaction without information flows. *The Review of Economic Studies*, 63(3):491–519.
- Grosskopf, B. and Sarin, R. (2010). Is reputation good or bad? an experiment. *American Economic Review*, 100(5):2187–2204.
- Holmström, B. (1999). Managerial incentive problems: A dynamic perspective. *The review of Economic studies*, 66(1):169–182.
- Holt, C. A. and Laury, S. K. (2002). Risk aversion and incentive effects. *American economic review*, 92(5):1644–1655.
- Kamijo, Y., Ozono, H., and Shimizu, K. (2016). Overcoming coordination failure using a mechanism based on gradualism and endogeneity. *Experimental Economics*, 19(1):202–217.
- Kartal, M., Müller, W., and Tremewan, J. (2021). Building trust: The costs and benefits of gradualism. *Games and Economic Behavior*, 130:258–275.
- Kranton, R. E. (1996). The formation of cooperative relationships. *The Journal of Law, Economics, and Organization*, 12(1):214–233.
- Kreps, D. M. and Wilson, R. (1982). Reputation and imperfect information. *Journal of economic theory*, 27(2):253–279.

- Kurzban, R., Rigdon, M. L., and Wilson, B. J. (2008). Incremental approaches to establishing trust. *Experimental Economics*, 11(4):370–389.
- Milgrom, P. and Roberts, J. (1982a). Predation, reputation, and entry deterrence. *Journal of economic theory*, 27(2):280–312.
- Milgrom, P. and Roberts, J. (1982b). Predation, reputation, and entry deterrence. *Journal of economic theory*, 27(2):280–312.
- Morris, S. (2001). Political correctness. *Journal of political Economy*, 109(2):231–265.
- Neral, J. and Ochs, J. (1992). The sequential equilibrium theory of reputation building: A further test. *Econometrica: Journal of the Econometric Society*, pages 1151–1169.
- Schelling, T. (1960). The strategy of conflict.
- Schmidt, K. M. (1993). Reputation and equilibrium characterization in repeated games with conflicting interests. *Econometrica: Journal of the Econometric Society*, pages 325–351.
- Sobel, J. (1985). A theory of credibility. *The Review of Economic Studies*, 52(4):557–573.
- Watson, J. (1999). Starting small and renegotiation. *Journal of economic Theory*, 85(1):52–90.
- Watson, J. (2002a). Starting small and commitment. *Games and Economic Behavior*, 38(1):176–199.
- Watson, J. (2002b). Starting small and commitment. *Games and Economic Behavior*, 38(1):176–199.
- Weber, R. A. (2006). Managing growth to achieve efficient coordination in large groups. *American Economic Review*, 96(1):114–126.

Ye, M., Zheng, J., Nikolov, P., and Asher, S. (2020). One step at a time: Does gradualism build coordination? *Management Science*, 66(1):113–129.



Appendix A

CHAPTER 1 PROOFS

Proof. **Lemma 1.1**

I first show that there are no reputation concerns in the last three periods in any equilibrium, i.e., both types of agents play the stage-game ideal actions (the action matching the state of world for the good agent, and action 1 for the bad agent). Suppose in an equilibrium, the agent plays a different strategy. The maximum payoff she can get in the last three periods is $k - 1 + 2k$. However, if she deviates to the stage-game ideal action, she gets a payoff weakly greater than k . $k < \frac{1}{2}$ implies $k > k - 1 + 2k$. Thus, in any equilibrium, $\mu_i^0 = 1$, $\mu_i^1 = 0$, $\nu_i = 0$ for $i \in \{n-2, n-1, n\}$. It in turn implies by Bayes' rule, $\rho_{i+1}(0) \geq \rho_{i+1}(1)$ for $i \in \{n-2, n-1\}$.

Consider period $n-3$. As there are no reputation concerns in the last three periods, agent's payoff is weakly increasing in ρ_{n-2} . Suppose for contradiction that $\rho_{n-2}(1) > \rho_{n-2}(0)$. Realize that this implies $\nu_{n-3} = 0$. If the good agent plays action 0 with a positive probability, then $\rho_{n-2}(0) = 1$, contradicting $\rho_{n-2}(1) > \rho_{n-2}(0)$. Thus, $\rho_{n-2}(1) > \rho_{n-2}(0)$ can only hold if $\mu_{n-3}^0 = 0$, $\mu_{n-3}^1 = 0$, i.e. both agents play action 1 with probability 1.

Sequential rationality of the agent implies that this can only be an equilibrium strategy if after observing $a_{n-3} = 1$, the agent is hired for at least two more periods. Suppose this is not the case. The highest payoff the agent can get is when the agent is hired in period $n-2$, but not hired in period $n-1$ if $a_{n-2} = 1$ is observed. Consider $\theta_{n-3} = 0$. Good agent gets payoff of $k - 1 + k + \frac{1}{2}2k$ if she plays $a_{n-3} = 1$. However, deviating to action 0 yields payoff k , which is strictly greater. Hence, $a_{n-3} = 1$ given $\theta_{n-3} = 0$ can not be an equilibrium strategy.

Let $m \geq 2$ as the number of periods for which the agent is hired following

$a_{n-3} = 1$ as long as no action 0 is observed afterwards. Then, the good agent's equilibrium payoff starting in period $n - 3$:

$$V_{n-3}^g = \left(k - \frac{1}{2}\right) + mk + \left(1 - \frac{1}{2^m}\right)(3 - m)k \quad (\text{A.1})$$

Decomposing A.1:

- $\left(k - \frac{1}{2}\right)$: Period $n - 3$ stage-game payoff.
- mk : As the agent is hired for m periods for sure (after any observation), she gets payoff of k in m periods.
- $\left(1 - \frac{1}{2^m}\right)(3 - m)k$: $a_i = 0$ is observed in any of the m periods with the probability $\left(1 - \frac{1}{2^m}\right)$. If that is the case, the agent is also hired for the remaining periods, $3 - m$ periods, and gets payoff of k in each period.

Let $\mu_i^0, \mu_i^1 = 0, \nu_i = 0$ for $i \in \{n - 3, n - 2, n - 1\}$. Under this strategies, the agent is hired for at least m periods starting in period $n - 3$ even if no action 0 is observed. The reasoning is as following. If the agent is hired in period i when the reputation is ρ' , then the agent is also hired in period $i - 1$ with the same reputation level. As the number of interaction periods increase, the principal hires the agent for weakly lower reputation levels. Under these strategies, the good agent gets the payoff at least:

$$V_{n-3}^g \geq mk + \left(1 - \frac{1}{2^m}\right)(4 - m)k \quad (\text{A.2})$$

Now I will show that this strategies constitute an equilibrium. The only profitable deviation for both agent is to deviate to action 0 in period $n - 3$. Such deviation yields a payoff of $k - 1 + 3k$ to both agent, which is strictly less than A.2, since $m \geq 2$ and $k < \frac{1}{2}$. Similarly, the bad agent does not have a profitable deviation. She gets payoff at least mk under the proposed strategy profile, which is strictly greater than her payoff from deviating to action 0: $k - 1 + 3k$.

Since A.2 is strictly higher than A.1, both type playing action 1 can not be a part of an agent-optimal equilibrium. This prove that, in the agent-optimal equilibrium $\rho_{i+1}(0) \geq \rho_{i+1}(1)$ for $i \in \{n-3, n-2, n-1\}$.

Next, I will show that the agent's payoff is weakly increasing in reputation in period $n-3$ in an agent-optimal equilibrium. Define ρ'_{n-2} such that $\lambda_{n-2} = 1$ if and only if $\rho_{n-2} \geq \rho'_{n-2}$. If ρ_{n-3} is so that $\frac{\rho_{n-3}}{2 - \rho_{n-3}} \geq \rho'_{n-2}$, then in any agent-optimal equilibrium, there is no reputation concern. More formally, $\mu_i^0 = \mu_i^1 = \nu_i = 0$ for $i \in \{n-3, n-2, n-1, n\}$. In this equilibrium, good agent gets the payoff weakly greater than $2k + \frac{3}{4}2k$. The higher ρ_{n-3} , the agent is hired for weakly more periods, hence gets a weakly higher payoff. This is the highest payoff agent can get in any equilibrium. In any equilibrium with reputation concerns, the good agent can get at most $k - 1 + 3k$, which is strictly less. Therefore, in any agent-optimal equilibrium, no reputation concern must be true.

Now consider $\rho_{n-3} < \frac{2\rho'_{n-2}}{1 + \rho'_{n-2}}$. No reputation can not be equilibrium strategy in this case because $\rho_{n-2}(0) = 1$ and the bad agent will deviate to play action 0. Thus, there will be reputation concerns in this case and the maximum payoff the agent can get is given by $k - 1 + 3k$. Observe that it is strictly less than the payoff of the agent when $\rho_{n-3} \geq \frac{2\rho'_{n-2}}{1 + \rho'_{n-2}}$. Also, I showed that the payoff is weakly increasing in reputation for $\rho_{n-3} \geq \frac{2\rho'_{n-2}}{1 + \rho'_{n-2}}$. Thus, the agent's payoff is weakly increasing in reputation in period $n-3$.

I finalize the proof with induction method. Consider period i such that $\rho_{j+1}(0) \geq \rho_{j+1}(1)$ and $\mu_j^0 = 1$ for all $j > i$. I will show that $\rho_{j+1}(0) \geq \rho_{j+1}(1)$ implies $\rho_{i+1}(0) \geq \rho_{i+1}(1)$.

Suppose $\rho_{i+1}(1) > \rho_{i+1}(0)$. Following the same reasoning above, it can only hold if both agent play action 1. There are 2 cases to consider.

Case 1

No reputation concern in any following period ($\{i, i+1, \dots, n\}$) is an equilibrium. Suppose ρ_i is such that the agent is hired for t periods if no action 0 is observed.

If $tk \geq k - 1 + (n - i)k$, then no agent has a profitable deviation, hence it is an equilibrium.

$\rho_{i+1}(1) > \rho_{i+1}(0)$ can only hold if both agents play action 1 with probability 1. Following the same steps above, for this to be a part of an equilibrium, ρ_i must be above certain level so that the agent is hired for $m \geq 1$ periods. Hence,

$$V_i^g = \left(k - \frac{1}{2}\right) + mk + \left(1 - \frac{1}{2^m}\right)(n - i - m)k \quad (\text{A.3})$$

On the other hand, in the equilibrium where there are no reputation concerns, the agent gets:

$$V_i^g \geq mk + \left(1 - \frac{1}{2^m}\right)(n - i - m)k \quad (\text{A.4})$$

A.4 is strictly higher than A.3, hence the no reputation yields a higher payoff to the agent.

Case 2

No reputation concern in any following period is not an equilibrium. More formally, consider ρ_i is such that the agent is hired for t periods if no action 0 is observed. If $tk < k - 1 + (n - i)k$, then the agent will deviate to play action 0 in period i .

Again $\rho_{i+1}(1) > \rho_{i+1}(0)$ can only hold if both agents play action 1 with probability 1. Following the same steps above, for this to be part of an equilibrium, ρ_i must be above certain level so that the agent is hired for $m \geq 1$ periods. The difference is that there will be reputation concerns in the period $i + 1$ (as $\mu_{i+1}^0 = 1$, $\nu_{i+1} = 0$ can not be an equilibrium strategy because it implies $\rho_{i+2}(0) = 1$). The maximum payoff the good agent can get in this formation:

$$V_i^g = \left(k - \frac{1}{2}\right) + \left(k - \frac{1}{2}\right) + (n - i - m)k \quad (\text{A.5})$$

On the other hand, consider the following equilibrium. The agent may screen her type in period i rather than in period $i + 1$:

$$V_i^g \geq \left(k - \frac{1}{2}\right) + (n - i - m)k \quad (\text{A.6})$$

A.6 is strictly higher than A.5, hence screening earlier yields a higher payoff to the agent. □

Proof. **Lemma 1.2**

Suppose for contradiction that there is an equilibrium where $\mu_i^0 < 1$ for some i . From Lemma 1.1, $\rho_{i+1}(0) > \rho_{i+1}(1)$ and $V_{i+1}^g(\cdot | a_i = 0) \geq V_{i+1}^g(\cdot | a_i = 1)$. Hence, the agent will deviate to $a_i = 0$ and increase both U_i^g and V_{i+1}^g . □

Proof. **Lemma 1.3**

From Lemma 1.1, $\rho_{i+1}(0) > \rho_{i+1}(1)$ implies $\nu_i < \frac{1 + \mu_i^1}{2}$. □

Proof. **Proposition 1.1**

Suppose for contradiction that there is an equilibrium where $\mu_i > 0$ for some i . The maximum continuation payoff the agent can get from playing action 0 is

$$k - 1 + (n - i)k$$

$k < \frac{1}{n-1}$ implies the agent will deviate to play action 1 and get the payoff of k . □

Proof. **Proposition 1.2**

(a) If the principal hires the agent in period i , then the expected payoff of the principal:

$$\gamma_i \left(\frac{1}{2} \rho_i k + \frac{1}{2} \rho_i (\mu_i (k - 1) + (1 - \mu_i) k) + (1 - \rho_i) \left(k - \frac{1}{2} \right) \right)$$

Which is negative for all values of μ_i and ν_i if $\rho_i < 1 - 2k$. Thus, the agent is not hired in the period i . Hence, the agent's reputation is not updated, and she will not be hired for all the remaining periods. The agent is never hired if $\rho_1 < 1 - 2k$.

(b) It is seen from the expected payoff in (a) that the expected payoff of the principal is non-negative if $\mu_i \leq \frac{2k + \rho_i - 1}{\rho_i}$. Hence the principal hires the agent only if the good agent plays action 0 with a sufficiently low probability. $\square$

Proof. **Proposition 1.3**

First, observe that $\mu_i = 0$ & $\nu_i = 0$ for all i is an equilibrium strategy for $\rho_1 \geq \frac{(1 - 2k)2^{n-2}}{(1 - 2k)2^{n-2} + k}$. By Bayes' rule, $\rho_n(1) \geq 1 - 2k$ after any history H_n and the agent is hired in every period irrespective of the history. The agent has no profitable deviation as deviation to play action 0 in any period i leads to the continuation payoff of $(n - i + 1)k - 1$ which is less than $(n - i + 1)k$.

The next thing is to show agent-optimality. The agent gets the payoff of k , which is the highest possible payoff she can get. In any equilibrium where the $\mu_i > 0$ ($\nu_i > 0$), the good (bad) agent gets a payoff less than k . $\square$

Proof. **Lemma 1.4**

Consider period t . Following the reasoning in Proposition 1.1, $n < \frac{1}{k} + t$ implies the agent has no reputation concern in all the remaining periods. Thus, $\mu_i = \nu_i = 0$ for $i \in t, t + 1, \dots, n$. The agent is hired in period t if $\rho_t \geq 1 - 2k$. By Bayes' rule, $\rho_1 \geq P(t)$ implies $\rho_t \geq 1 - 2k$ if $\mu_j = 0 \forall j < t$.

Final step is to show $\mu_j = 0 \forall j < t$. Consider period t and a history h such that $\rho_t \geq 1 - 2k$. Sequential rationality constraint of the good agent implies $\mu_t = 0$ irrespective of the bad agent's strategy. Agent gets a minimum payoff of k , and she can get maximum payoff of $k - 1 + (n - t)k$ by playing action 0 with a positive probability when $\theta_t = 1$. $n < \frac{1}{k} + t$ implies the agent has no incentive to set $\mu_t > 0$. The same reasoning applies to the bad agent, hence $\nu_t = 0$. Continuing with backward induction, we observe that good agent has no incentive to set $\mu_j > 0$ for any $j < t$. Bad agent follows. $\square$

Proof. **Lemma 1.5**

From Proposition 1.2, the agent is not hired in period i if $\mu_i > \mu'_i = \frac{2k + \rho_i - 1}{\rho_i}$. Hence $\mu_i < 1$ in any period where the agent is hired. $\mu_i < 1$ implies $\nu_i < 1$ from Lemma 1.3. I will show that in any equilibrium, if $\rho_{i+1}(0) < P(n - i) = \frac{(1 - 2k)2^{n-i-2}}{(1 - 2k)2^{n-i-2} + k}$, then $\nu_i > 0$ implies $\mu_i = 1$, i.e. the good agent has more reputational incentives compared to the bad agent.

Consider the last period i in which the bad agent has a reputational incentive in an equilibrium. Following the reasoning in Proposition 1.1, there is no reputational incentives in the last three periods, hence there is such period $i \leq n - 2$. In such period, the reputation level must be close enough to $1 - 2k$ such that it is a critical period and the agent is not hired for the remaining periods if action 1 is observed.. For contradiction, consider period i such that the bad agent is hired for at least 1 periods after action 1 is observed and he has no reputational concern for the remaining periods. Define the number of periods m and l as following. If $a_i = 0$ and action 1 is observed for the remaining periods, the agent is hired for m periods. As the bad agent will not play action 0 with a positive probability, she will be hired for exactly m periods after $a_i = 0$. On the other hand, if $a_i = 1$ is observed and no action 0 is observed in the succeeding periods, then the agent is hired for l periods. By the same reasoning, the bad agent is hired for exactly l periods after $a_i = 1$. Note that $n - i \geq m > l$. I will show that if period i is the last period where the bad agent has reputational incentives, then $l = 0$.

The sequential rationality constraint of the bad agent for $\nu_i > 0$ implies:

$$k - 1 + mk = k + lk \iff mk - 1 = lk$$

Consider period $i + 1$ such that $a_i = 1$ is observed. I assumed $\nu_{i+1} = 0$, hence $\rho_{i+2}(0) = 1$. If the bad agent plays $a_i = 1$, then she will be hired for $l - 1$ periods. On the other hand, if she deviates to play $a_{i+1} = 0$, then she will get a payoff of $(n - i)k - 1$. If $n - i > m$, then the bad agent will deviate to play action 0. If $n - i = m$, then the bad agent continues to hold reputational incentives. Thus, in the last period where the bad agent holds reputational incentives, the agent is not

hired after action 1.

Now, consider such a critical period where the reputation is sufficiently low, and the bad agent has reputational incentives. $l = 0$ and suppose that $m < n - i$. This assumption implies that $\rho_{i+1}(0) < P(n - i)$. By Bayes' rule, if action 1 is observed in the following periods, then $\rho_n < 1 - 2k \Leftrightarrow m < n - i$. From the sequential rationality constraint of the bad agent, $mk = 1$. Consider the sequential rationality constraint of the good agent. Conditioning on $\theta_i = 1$, the expected continuation payoff of the good agent from playing action 0 is as following:

$$k - 1 + \frac{1}{2^m}mk + \left(1 - \frac{1}{2^m}\right)(n - i)k$$

$\frac{1}{2^m}$ is the probability with which $\theta = 1$ for the next m periods. Then the agent is not hired after the period $m + i$. With the probability $(1 - \frac{1}{2^m})$, $\theta = 0$ for some for some period, hence the good agent fully reveals her type and is hired for all the remaining periods, i.e. $n - i$ periods. $mk = 1$ and $m < n - i$ implies the expected payoff above is higher than k . Thus, the best response to $\nu_i > 0$ is $\mu_i = 1$ and the agent is not hired in period i . Note that,

$$k - 1 + \frac{1}{2^m}mk + \left(1 - \frac{1}{2^m}\right)(n - i)k \leq k$$

Which can be simplified as follows:

$$\left(1 - \frac{1}{2^m}\right)(n - i)k \leq 1 - \frac{1}{2^m}$$

Where I used the fact that $mk = 1$. Thus I need $(n - i)k \leq 1$ to hold for $\mu_i < 1$. $m < n - i$ implies $(n - i)k > 1$, contradiction. Hence, the SR constraint of the good agent holds only if $n - i = m$, i.e. $\mu_i \leq \mu'_i$ can hold.

I showed that in the last period in which the bad agent holds reputational concerns, the good agent has strictly higher reputational concerns if $\rho_{i+1}(0) < P(n - i)$. Only if $\rho_{i+1}(0) \geq P(n - i)$, both types are indifferent between the two actions. Hence, the bad agent holding reputational incentives in period i implies $(n - i)k = 1$.

Next consider the period $j < i$ such that $\nu_j \geq 0$ and $\nu_i \geq 0$. Hence, $(n - i)k = 1$. Note that in the periods between j and i , as $\nu < 1$, the bad agent gets the payoff of k in each period. Suppose $\rho_{j+1}(0) < P(n - j)$, i.e. the agent is hired for m period after action 0 such that $m < n - j$. The sequential rationality constraint of the bad agent for $\nu_j \geq 0$ implies:

$$k - 1 + mk = k + (i - j)k \iff mk - 1 = (i - j)k$$

$(n - i)k = 1$ implies $ik = nk - 1$. The right-hand side of the equation above becomes $(n - j)k - 1$. $m < n - j$ implies $mk - 1 < (n - j)k - 1$. Hence, $mk - 1 = (i - j)k$ can not hold for any $m < n - j$. Thus, there is no equilibrium in which the agent is hired in period i and ν_i is such that $\rho_{i+1}(0) < P(n - i)$.

□

Proof. Proposition 1.4

From Lemma 1.4, if $n < \frac{1}{k} + t$, then the agent has no reputational concern. Hence, the agent is first hired in the first period and in the history in which no action 0 is observed, the agent is hired for first t periods.

Consider $n > \frac{1}{k} + t$ and $P(t) \leq \rho_1 < P(t + 1)$. I will show that agent will not be hired in any period with reputation concerns. Note that, if hiring starts in period $i < n - t + 1 - \frac{1}{k}$, then $\mu_j = \nu_j = 0$ for all the periods agent is hired can not be an equilibrium strategy. Consider the bad agent. If she does not take the costly reputation action, then she will be hired for t periods. She gets payoff of tk . On the other hand, if she deviates to play action 0 in any period in $i, i + 1, \dots, i + t - 1$, she gets payoff of $(n - i)k + k - 1$. It is profitable to deviate as $i < n - t + 1 - \frac{1}{k}$. Hence, if hiring starts in period $i < n - t + 1 - \frac{1}{k}$, then there will be reputation concerns.

From Proposition 1.2, the agent is not hired in any period i where $\mu_i \leq \frac{2k + \rho_i - 1}{\rho_i}$. Hence, if $\mu_i = 1$, i.e. the good agent has strict reputation incentives, then the agent is not hired. I will show that if the bad agent has reputation concerns in any period,

then the good agent has strict reputation concerns, i.e. $\nu_i > 0 \implies \mu_i = 1$.

From Lemma 1.5, $\nu_i > 0 \implies \mu_i = 1$ unless $\rho_{i+1}(0) \geq P(n-i)$. Hence, for the agent to be hired in any period i with reputation concerns, the following must hold.

$$\rho_{i+1}(a_i = 0) = \frac{\rho_i + \rho_i \mu_i}{\rho_i + \rho_i \mu_i + 2(1 - \rho_i) \nu_i} \geq \frac{(1 - 2k)2^{n-i-2}}{(1 - 2k)2^{n-i-2} + k}$$

Which implies

$$\nu_i \leq \frac{k(1 + \mu_i) \rho_i}{2^{n-i-1}(1 - 2k)(1 - \rho_i)}$$

Hence, by Bayes' rule, $\rho_i < P(t+1)$ implies $\rho_{i+1}(1) < P(t)$. It means, the agent is hired for at most t periods after $a_i = 1$. Then the bad agent deviates to play action 0 with probability 1, contradiction. Hence, there is no such period i where $\nu_i > 0$ and $\rho_{i+1}(0) \geq P(n-i)$. The same reasoning applies to any period $j \in i, i+1, \dots, i+t-1$. Thus, if the bad agent has reputation concerns in any period j , then $\mu_j = 1$. The agent is not hired in such period. Therefore, the agent is not hired in period $i < n - t + 1 - \frac{1}{k}$.

□

Proof. **Lemma 1.6**

In any period, the principal hires the agent if the expected continuation payoff is non-negative. Consider period j so that ρ_j is in the lowest range for which the agent is hired. It implies that if action 1 is observed, then $\lambda_{j+1}(1) = 0$. Thus, the expected continuation payoff of the principal in period j is given by:

$$V_j^P = \left(k - \frac{1}{2} + \frac{1}{2} \rho_j \right) + \frac{1}{2} \rho_j (n-j)k$$

Where $\frac{1}{2} \rho_j$ denotes the probability of observing $a_j = 0$. This payoff is non-negative

for $\rho_j \geq \frac{1-2k}{1+(n-j)k}$, hence $\frac{1-2k}{1+(n-j)k}$ is the lowest reputation level for which the agent is hired. □

Proof. Proposition 1.5

1. First, observe that $\mu_i = 0 \& \nu_i = 0$ for all i is an equilibrium strategy for

$\rho_1 \geq \frac{(1-2k)2^{n-2}}{(1-2k)2^{n-2} + k}$. By Bayes' rule, $\rho_n(1) \geq 1-2k$ and the agent is hired in every period irrespective of the history. The agent has no profitable deviation as deviation to play action 0 in any period i leads to the continuation payoff of $(n-i+1)k-1$ which is less than $(n-i+1)k$.

The next thing is to show agent-optimality. The agent gets the payoff of k , which is the highest possible payoff she can get. In any equilibrium where the $\mu_i > 0$ ($\nu_i > 0$), the good (bad) agent gets a payoff less than k .

2. Consider period t . Following the reasoning in Proposition 1.1, $n < \frac{1}{k} + t$ implies the agent has no reputation concern in all the remaining periods. Thus, $\mu_i = \nu_i = 0$ for $i \in t, t+1, \dots, n$. From Lemma 1.6, the agent is hired in period t if $\rho_t \geq \frac{1-2k}{1+(n-t)k}$. By Bayes' rule, $\rho_1 \geq R(t, n)$ implies $\rho_t \geq \frac{1-2k}{1+(n-t)k}$ if $\mu_j = 0 \forall j < t$.

Final step is to show $\mu_j = 0 \forall j < t$. Consider period t and a history h such that $\rho_t \geq \frac{1-2k}{1+(n-t)k}$. Sequential rationality constraint of the good agent implies $\mu_t = 0$ irrespective of the bad agent's strategy. Agent gets a minimum payoff of k , and she can get maximum payoff of $k-1+(n-t)k$ by playing action 0 with a positive probability when $\theta_t = 1$. $n < \frac{1}{k} + t$ implies the agent has no incentive to set $\mu_t > 0$. The same reasoning applies to the bad agent, hence $\nu_t = 0$. Continuing with backward induction, I observe that good agent has no incentive to set $\mu_j > 0$ for any $j < t$. Bad agent follows.

3. Following Proposition 1.1, $n < \frac{1}{k} + 1$ implies the agent has no reputation concern in any period. From Lemma 1.6, the principal's expected payoff in the first period is given by:

$$V_1^P = \left(k - \frac{1}{2} + \frac{1}{2}\rho_1 \right) + \frac{1}{2}\rho_1(n-1)k$$

Which is non-negative if $\rho_1 \geq R(1, n)$. Hence the agent is not hired for the lower initial reputation levels.

This completes the proof. $\square$

Proof. **Proposition 1.6**

Directly follows Proposition 1.5. $\square$

Proof. **Proposition 1.7**

I will prove in three steps.

Step 1.

Claim: $\nu_1 \in (0, 1)$ must be in any PBE.

$\nu_1 = 0$ can not be an equilibrium strategy. Suppose for contradiction that $\nu_1 = 0$. Then $\rho_2(0) = 1$ by Bayes' rule. Hence, $V_2^b(.|a_1 = 0) = k + k + \dots + k$. $n > \frac{1}{k} + 1$ implies $k - 1 + k + k + \dots + k > k$, hence the sequential rationality implies the agent will deviate to play action 0.

$\nu_1 = 1$ can not be an equilibrium strategy. Suppose $\nu_1 = 1$. Then $\mu_1 = 1$ and $\rho_2(0) = \rho_1$. Consider period 2. Two cases are possible.

Case 1. $\lambda_2(0) = \lambda_2(1) = 0$. But then, the bad agent will deviate to action 1 in period 1.

Case 2. $\lambda_2(0) = 1$ and $\lambda_2(1) = 0$. $\frac{1}{k} + 2 \geq n$ implies the agent will no more have reputation concerns. Hence, $V_2^b = k$. But then, the bad agent will deviate to action 1 in period 1. Thus, $\nu_1 \in (0, 1)$ must be in any PBE.

Step 2.

Claim: $\mu_1 = 1$.

Note that the agent is not hired after action 1 and hired for some periods if action 0 is observed in the period 1. It implies $V_1^g(.|a_1 = 0) \geq V_1^b(.|a_1 = 0)$. The reasoning is that for each period that the agent is hired, there is $1/2$ probability for which the good agent can invest in her reputation costless - when the state of world is 0. Thus, the good agent has a higher reputational incentive compared to the bad agent. Suppose $V_1^g(.|a_1 = 0) = V_1^b(.|a_1 = 0)$. It only holds if the agent is hired in all periods in the future if action 0 is observed. This is the only case where both type of agent has the same sequential rationality constraint. However, $n > \frac{1}{k} + 1$ implies

that bad agent will deviate to play action 0 with probability 1. But then $\rho_2(0) = \rho_1$, contradiction. Thus, $V_1^g(\cdot|a_1 = 0) \neq V_1^b(\cdot|a_1 = 0)$. $\nu_1 \in (0, 1)$ implies $\mu_1 = 1$ in the agent-optimal equilibrium.

Step 3.

Next, I will define the principal optimality. Note that the expected continuation payoff of the principal is decreasing in ν_1 as $\rho_2(0)$ is decreasing in ν_1 . In the principal-optimal equilibrium ν_1 is equal to the lowest value satisfying the equilibrium conditions. The equilibrium with the lowest ν_1 , hence the highest $\rho_2(0)$ is the principal-optimal equilibrium. I will show that the equilibrium with the lowest ν_1 is also agent-optimal. The highest $\rho_2(0)$ is such that the principal will hire the agent for all the remaining periods irrespective the observed action in each period. But following the reasoning above, the bad agent would deviate to play action 1 with probability 1. The second best alternative is that in the history h_n in which $a_i = 1 \forall i \in N/\{1, n\}, \rho_n(1) = 1 - 2k$. In words, if the action 1 is observed in each period, then the principal is indifferent between hiring the agent or not in period n . Bayes' rule implies if $\rho_2 = \frac{(1 - 2k)2^{n-3}}{(1 - 2k)2^{n-3} + k}$, then $\rho_n(a_{n-1} = 1|a_j = 1, j \in \{2, \dots, n-2\}) = 1 - 2k$. Hence, ν_1 will be so that

$$\rho_2(0) = \frac{\rho_1}{\rho_1 + (1 - \rho_1)\nu_1} = \frac{(1 - 2k)2^{n-3}}{(1 - 2k)2^{n-3} + k}$$

$$\Rightarrow \nu_1 = \frac{\rho_1 k}{2^{n-3}(1 - 2k)(1 - \rho_1)}$$

Sequential rationality constraint of the bad agent becomes:

$$k = k - 1 + k + k + \dots + k + \alpha k$$

Where α stands for $\lambda_n(1)$ given action 1 is observed in each period except period 1.

It implies $\alpha = \frac{1}{k} - (n - 2)$. Note that, as $\mu_1 = 1$ in any equilibrium, the equilibrium with the highest $\rho_2(0)$ is also optimal for the good agent. With same reasoning, as $\nu_1 \in (0, 1)$ in the agent-optimal equilibrium, thus $V_1^b = k$. Hence, the above defined equilibrium is also optimal for the bad agent. $\square$

Proof. Proposition 1.8

In period i , $i > n - \frac{1}{k}$ implies the agent has no reputational concern in any equilibrium.

Consider $i = n - \frac{1}{k}$. Note that it is an exceptional case in which $\frac{1}{k}$ is an integer. Suppose ρ_i is in the smallest range in period i so that the agent is hired in period i but if action 1 is observed, he is not hired thereafter. The sequential rationality constraint of the agent implies she is indifferent between two actions in this case:

$$k = k - 1 + (n - i)k$$

Note that, $i = n - \frac{1}{k}$ implies $i + j > n - \frac{1}{k}$ for any $j \geq 1$. Hence, the agent will not have reputational concern in any period $i + j$. The agent is hired if $\rho_{i+1} \geq$

$\frac{1 - 2k}{1 + (n - i - 1)k}$. If the reputation is in the smallest range in period i , it implies that $\rho_{i+1}(1) < \frac{1 - 2k}{1 + (n - i - 1)k}$. Consider that $\rho_{i+1}(1) \geq \frac{\rho_i}{2 - \rho_i}$, which holds with equality if $\mu_i = 0$. Hence, lowest reputation level in period i is equivalent to:

$$\frac{\rho_i}{2 - \rho_i} < \frac{1 - 2k}{1 + (n - i - 1)k}$$

$$\implies \rho_i < \frac{2 - 4k}{2(n - i - 2)k + 1} = \frac{2 - 4k}{3 - 4k}$$

The last equality follows from $k = \frac{1}{n - i}$.

$$\begin{aligned} V_i^P &= \left(\frac{1}{2}\rho_i k + \frac{1}{2}\rho_i \mu_i (k - 1) + \frac{1}{2}\rho_i (1 - \mu_i)k + (1 - \rho_i) \left(k - \frac{1}{2} \right) \right) + \\ &+ \left(\frac{1}{2}\rho_i + \frac{1}{2}\rho_i \mu_i \right) (n - i)k \end{aligned}$$

Which is independent of μ_i for $k = \frac{1}{n - i}$. It is reasonable as the good agent and the principals share the same stage game preferences. Hence, if the agent is

indifferent between two actions for certain parameter values, it is natural to expect the principal's payoff to be independent of the agent's mixing probability. Thus, the

principal hires the agent if $\rho_i \geq \frac{1-2k}{1+(n-i)k} = 1 - \frac{1}{2}k$.

Thus, in any period i , if $k = \frac{1}{n-i}$, then the lowest reputation range for which the agent is hired is as follows:

$$1 - \frac{1}{2}k \leq \rho_i < \frac{2-4k}{3-4k}$$

I showed that the agent is indifferent between two actions if $k = \frac{1}{n-i}$ in period i if $\rho_i < \frac{2-4k}{3-4k}$. Next, I will show that, the agent prefers action 1 for the higher levels

of the reputation. I will show that, for any $\rho_i \geq \frac{2-4k}{3-4k}$, the agent does not play action 0 with a positive probability in the agent-optimal equilibrium. Consider the following equilibrium. $\rho_i \geq \frac{2-4k}{3-4k}$, and $\mu_i > 0$. $k = \frac{1}{n-i}$ implies there is such an equilibrium where μ_i high enough so that the agent is not hired after action 1. In this

equilibrium, the agent gets the expected continuation payoff of $(n-i+1)k - \frac{1}{2} = \frac{1}{2} + k$ in period i . With probability $\frac{1}{2}$, $\theta_i = 1$ and thus there is $-\frac{1}{2}$ in the payoff. Now

consider the equilibrium with $\rho_i \geq \frac{2-4k}{3-4k}$ and $\mu_i = 0$. Following the analysis above,

$k = \frac{1}{n-i}$ implies the agent does not have a reputational concern in the remaining periods. Hence, in this equilibrium, the agent gets the expected continuation payoff of $k + \frac{1}{2}(n-i)k + \frac{1}{2} \left(k + \frac{1}{2}(n-i-1)k \right) = \frac{5}{4}k + \frac{3}{4}$ in period i . It is clear that the agent gets a higher expected continuation payoff in this equilibrium. Hence, in the agent optimal equilibrium $\mu_i = 0$ for $\rho_i \geq \frac{2-4k}{3-4k}$.

1. Consider $i < n - \frac{1}{k}$:

I will first show that the agent postpones screening her type as long as her reputation is not so low in the agent-optimal equilibrium. Define ρ'_i such that $\rho_i \geq \rho'_i$ implies $\lambda_{i+1}(1) = 1$. In words, for given reputation levels, the agent is hired in period i , and even if action 1 is observed, she is hired for at least 1 more period(s). Consider the following equilibrium. $\rho_i \geq \rho'_i$, and $\mu_i > 0$. There is such $k > \frac{1}{n-i}$ for which the agent has a reputational incentive in an equilibrium. Without loss of generality, suppose $\mu_i = 1$. In this equilibrium, the agent gets the expected continuation payoff of $(n-i+1)k - \frac{1}{2}$ in period i . With probability $\frac{1}{2}, \theta_i = 1$ and thus we have $-\frac{1}{2}$ in the payoff. Now consider the equilibrium with $\rho_i \geq \rho'_i$ and $\mu_i = 0$. Suppose that the agent postpones screening by only periods. She get expected continuation payoff of $(n-i+1)k - \frac{1}{4}$. In this case, she needs to screen her type costly with only the probability of $\frac{1}{4}$, which is the probability of state being 1 in periods i and $i+1$. Hence, in the agent-optimal equilibrium, if the agent knows that she will be hired after action 1, then she will not play the costly action - action 0. She postpones screening as far as possible.

Next, I will justify the hiring decision of the principal. From above analysis, $i < n - \frac{1}{k}$ and sufficiently low reputation level implies the agent plays action 0 with probability 1, i.e. $\mu_i = 1$. Then:

$$V_i^P = \left(k - \frac{1}{2}\right) + \rho_i(n-i)k$$

Which is non-negative for $\rho_i \geq \frac{1-2k}{2(n-i)k}$. Hence,

$$\lambda_i = \begin{cases} 0, & \rho_i < \frac{1-2k}{2(n-i)k} \\ 1, & \rho_i \geq \frac{1-2k}{2(n-i)k} \end{cases}$$

The last thing is to show ρ'_i . ρ'_i depends on whether the agent has a reputation incentive in the next period or not. It, in turn, depends on the relation between

$i + 1$ and $n - \frac{1}{k}$. If $i + 1 \leq n - \frac{1}{k}$, it means the agent has a reputational concern in the next period if her reputation falls below a certain level. On the other hand, $i + 1 > n - \frac{1}{k}$ implies the agent will have no reputation concern for the remaining periods.

First consider $i + 1 > n - \frac{1}{k}$. It implies, in period $i + 1$, if the agent is hired, then she will have no more reputational concern - $\mu_{i+1} = 0$. Thus

$$V_{i+1}^P = \left(k - \frac{1}{2} + \frac{1}{2}\rho_{i+1} \right) + \frac{1}{2}\rho_{i+1}(n - i - 1)k$$

and $\lambda_{i+1} = 1$ if $\rho_{i+1} \geq \frac{1 - 2k}{1 + (n - i - 1)k}$.

It implies, if $\rho_i \geq \frac{2 - 4k}{2 + (n - i - 3)k}$ and $\mu_i = 0$, then the agent is hired for at least 1 more periods if action 1 is observed. Thus, the agent has no reputational concern.

Hence $\rho'_i = \frac{2 - 4k}{2 + (n - i - 3)k}$ if $i + 1 > n - \frac{1}{k}$.

Now consider $i + 1 \leq n - \frac{1}{k}$.

It implies, in period $i + 1$, if the agent is hired and her reputation is sufficiently small, then $\mu_{i+1} \geq 0$. Note that, if $i + 1 = n - \frac{1}{k}$, then $\mu_{i+1} \in [0, 1]$, otherwise $\mu_{i+1} = 1$. Thus

$$V_i^P = \left(k - \frac{1}{2} \right) + \rho_i(n - i - 1)k$$

and $\lambda_{i+1} = 1$ if $\rho_{i+1} \geq \frac{1 - 2k}{2(n - i - 1)k}$. It implies, if $\rho_i \geq \frac{2 - 4k}{2(n - i - 2)k + 1}$ and $\mu_i = 0$, then the agent is hired for at least 1 more periods if action 1 is observed.

Thus, the agent has no reputational concern. Hence $\rho'_i = \frac{2 - 4k}{2(n - i - 2)k + 1}$ if $i + 1 \leq$

$n - \frac{1}{k}$. Note that $i + 1 = n - \frac{1}{k}$ implies the agent will be indifferent in period $i + 1$

and $\frac{2-4k}{2(n-i-2)k+1} = \frac{2-4k}{2+(n-i-3)k} = \frac{2-4k}{3-4k}$. So, I have different ρ'_i values summarized below:

$$\rho'_i = \begin{cases} \frac{2-4k}{2(n-i-2)k+1}, & k \geq \frac{1}{n-i-1} \\ \frac{2-4k}{2+(n-i-3)k}, & k < \frac{1}{n-i-1} \end{cases}$$

□

Appendix B

CHAPTER 2 PROOFS

Proof. **Proposition 2.1**

First, note that in the last period (period n), $\mu_n = 0 = \nu_n$ as no type has a reputation incentive.

Take period i and consider the following: $\rho_i \geq \frac{(1-2k)2^{n-i-1}}{(1-2k)2^{n-i-1} + k}$, $\mu_i = 0 = \nu_i$ $\forall \delta_i$ and the choice of delta is arbitrary: $\delta_i \in [\underline{\delta}, \bar{\delta}]$ and $\lambda_{i+1}(0) = 1 = \lambda_{i+1}(1)$. Note that $\rho_i \geq \frac{(1-2k)2^{n-i-1}}{(1-2k)2^{n-i-1} + k}$ implies $\rho_n \geq 1 - 2k$ if $\mu_j = 0 = \nu_j$ for all $j > i$. Then the agent gets the maximum possible payoff as she is hired in each period and maximizes the stage-game payoff.

Now suppose there is an equilibrium in which $\mu_i > 0$ or $\nu_i > 0$ for some i where $\rho_i \geq \frac{(1-2k)2^{n-i-1}}{(1-2k)2^{n-i-1} + k}$. As $\gamma_i > 0$, the agent gets a continuation payoff strictly less than $(n-i+1)k$. Hence, the agent will deviate to set a lower δ_i and play the stage-game optimal action.

Thus, there will be no reputational concern in the agent-optimal equilibrium and the agent sets stakes arbitrarily. The choice of delta comes from the fact that $\mu_i = 0 = \nu_i$ for all $i \in N$ is an equilibrium strategy for any δ_i . As the agent is going to be hired for all the periods, her continuation payoff does not depend on the allocation of the relative stakes.

Note that, in any equilibrium, the bad agent will pool with the good agent when setting the δ_i . Suppose not. The bad agent sets $\delta_i = \delta'_i$ and the good agent sets $\delta_i = \delta''_i$. Bayes' rule implies $\rho_i(\delta_i|h_i) = 0$ for any $\delta_i \neq \delta''_i$ on the equilibrium path. Hence, $\lambda_i(\delta_i) = 0$ for any $\delta_i \neq \delta''_i$. The bad agent deviates to play $\delta_i = \delta''_i$. Hence, I call agent's decision of δ_i , I will refer to good agent. Moreover, I will show that the choice of equilibrium δ_i is also optimal for the bad agent even if I assume δ_i carries

no signal about the agent's type.

The agent will never set δ_i so that she has a reputation incentive. From the previous analyses, we know that $\mu_i \leq \mu'_i$ in any period i where the agent is hired. Hence, if $\mu_i > 0$, then $\mu_i \in (0, \mu'_i]$. This implies, the agent is indifferent between both actions in the period i . Note that, given $\theta_i = 1$, the highest expected continuation payoff agent can get is given by $V_i^g = (1 - \delta_i) \left(k - \frac{1}{2} \right) + \delta_i k$.

Now consider the following, $\delta_i = \underline{\delta}$ and $\mu_i = 0 = \nu_i$. The minimum continuation payoff of the agent is $(1 - \underline{\delta})k$ which is greater than $(1 - \delta_i) \left(k - \frac{1}{2} \right) + \delta_i k$ as $\delta_i \leq \bar{\delta} = 1 - \underline{\delta}$. Hence, as there is a profitable deviation, setting δ_i such that there is a reputational incentive can not be agent optimal equilibrium strategy.

Next I show that if $\mu_i = 0$, then the agent's expected continuation payoff is decreasing in δ_i , i.e. $\delta_i = \underline{\delta}$ is the optimal choice.

If $\mu_i = 0$, $V_i^g = (1 - \delta_i)k + \delta_i V_{i+1}^g$. Note that, $V_{i+1}^g < k$. $\rho_i \leq \frac{(1 - 2k)2^{n-i-1}}{(1 - 2k)2^{n-i-1} + k}$ implies that if no action 0 is observed for some periods $\{i, i + 1, \dots, j\}$, then $\rho_j < \frac{1 - 2k}{1 - k}$ and $\lambda_{j+1}(1) = 0$. $V_j^g = (1 - \delta_j)k + \delta_j V_{j+1}^g$. Since, $\lambda_{j+1} = 0$ with probability $\frac{1}{2}$, V_j^g is decreasing in δ_j and $\delta_j = \underline{\delta}$ is optimal for the agent. As $\underline{\delta} > 0$, $V_j^g < k$. As there is a positive probability that the agent will reach a node where she will get payoff less than k , $V_{i+1}^g < k$. Hence, V_i^g is decreasing in δ_i and the maximizing level is at $\delta_i = \underline{\delta}$.

As the bad agent pools with the good agent in the choice of delta, $\delta_i = \underline{\delta}$ constitutes the agent-optimal perfect Bayesian equilibrium.

Note that the results also applies to the setup where the bad agent is a commitment type. If the bad agent is a commitment type, the good agent still has two choices of delta in each period:

1. Set δ_i such that the agent reveals her type. In this case, the good agent will not set any delta higher than the value for which she is indifferent between two actions. The reasoning follows from the analysis above. If $\mu_i = 1$, then she will

not be hired. Hence, if δ_i is such that she is indifferent between two actions, she will prefer to deviate to a lower delta and play action 1. The comparison of the payoffs is the same as the comparison above.

2. Set δ_i such that the agent has no reputational incentive. For such deltas, $\delta_i = \underline{\delta}$ maximizes the expected continuation payoff of the agent. Hence, we have the same equilibrium when the bad agent is the commitment type.

As $\mu_i = 0$ for all i , the principal hires the agent as long as $\rho_i \geq 1 - 2k$.

□

Proof. **Proposition 2.2**

The agent has reputational incentives for high enough δ_i in the initial periods. First, I will show that the social planner (SP) will not set δ_i to a level for which the agent has strict reputational incentives, i.e. $\mu_i = 1$. In such period, the agent is not hired. As δ_i is bounded above by $\bar{\delta}$, $1 - \bar{\delta}$ proportion of the remaining decisions is lost in period i , Hence, the social planner will deviate to a lower δ_i so that the agent is hired.

Next, I will use induction method to show that the social planner will not set δ_i so that the agent strictly prefers to play action 1, i.e. $\mu_i = 0$. Consider period i such that $\rho_i < \frac{1 - 2k}{1 - k}$. Bayes' rule implies $\rho_{i+1}(1) < 1 - 2k$, hence $\lambda_{i+1}(1) = 0$. By

sequential rationality constraint of the agent, $\mu_i = 0$ if $\delta_i \leq \frac{1}{1 + k}$.

$$V_i^{SP} = (1 - \delta_i) \left(k - \frac{1}{2} + \frac{1}{2}\rho_i \right) + \delta_i \frac{1}{2}\rho_i k$$

Which is increasing in δ_i as $\rho_i < \frac{1 - 2k}{1 - k}$. Thus, the SP will deviate to $\delta_i = \frac{1}{1 + k}$. V_i^{SP} becomes

$$V_i^{SP} = \frac{k}{1 + k} \left(k - \frac{1}{2} + \rho_i \right)$$

Which is independent of μ_i .

Consider period $i - 1$ and suppose that δ_{i-1} is such that $\mu_{i-1} = 0$. By sequential rationality constraint of the agent, $\mu_{i-1} = 0$ if $\delta_{i-1} \leq \frac{2+2k}{2+3k}$. Note that $\rho_{i-1} < \frac{2-4k}{2-3k}$ implies $\rho_i < \frac{1-2k}{1-k}$ by Bayes rule.

$$V_{i-1}^{SP} = (1 - \delta_{i-1}) \left(k - \frac{1}{2} + \frac{1}{2}\rho_{i-1} \right) + \delta_{i-1} \frac{1}{2}\rho_{i-1}k + \\ + \delta_{i-1} \left(1 - \frac{1}{2}\rho_{i-1} \right) \frac{k}{1+k} \left(k - \frac{1}{2} + \rho_i(1) \right)$$

Where $\frac{1}{2}\rho_{i-1}$ denotes the $\text{prob}[a_i = 0]$ and $1 - \frac{1}{2}\rho_{i-1}$ denotes the $\text{prob}[a_i = 1]$ and $\rho_i(1) = \frac{\rho_{i-1}}{2 - \rho_{i-1}}$. $\rho_{i-1} < \frac{2-4k}{2-3k}$ implies Π_{i-1}^{SP} is increasing in δ_{i-1} , hence the payoff maximizing $\delta_{i-1} = \frac{2+2k}{2+3k}$.

Next consider period j such that δ_{j+1} is equal to δ_{j+1}^g where the agent is indifferent between both actions. Suppose δ_j is such that $\mu_j = 0$.

$$V_j^{SP} = (1 - \delta_j) \left(k - \frac{1}{2} + \frac{1}{2}\rho_j \right) + \delta_j \frac{1}{2}\rho_j k + \delta_j \left(1 - \frac{1}{2}\rho_j \right) V_{j+1}^{SP}(\cdot | a_j = 1)$$

$\delta_{j+1} = \delta_{j+1}^g$ implies that $V_{j+1}^{SP}(\cdot | a_j = 1)$ is increasing in δ_{j+1} for $\delta_{j+1} \leq \delta_{j+1}^g$. It implies that $V_{j+1}^{SP}(\cdot | a_j = 1) > k - \frac{1}{2} + \frac{1}{2}\rho_{j+1}(1)$ which is equal to $k - \frac{1}{2} + \frac{1}{2} \frac{\rho_j}{2 - \rho_j}$. It implies

$$\left(1 - \frac{1}{2}\rho_j \right) V_{j+1}^{SP}(\cdot | a_j = 1) > \left(k - \frac{1}{2} - \frac{1}{2}\rho_j k + \frac{1}{2}\rho_j \right)$$

Thus,

$$\delta_j \frac{1}{2} \rho_j k + \delta_j \left(1 - \frac{1}{2} \rho_j \right) V_{j+1}^{SP}(\cdot | a_j = 1) > \delta_j \left(k - \frac{1}{2} + \frac{1}{2} \rho_j \right)$$

Hence, V_j^{SP} is also increasing in δ_j and the maximizing $\delta_j = \delta_j^g$. The social planner allocates stakes so that the agent is indifferent between the two actions in any period.

Next, I will show by iteration that $\delta_i^g = \frac{2^{n-t+1-i}(1+2k) - 2k}{2^{n-t+1-i}(1+2k) - k}$.

In period $n - t + 1$, $1 - 2k \leq \rho_{n-t+1} < \frac{1 - 2k}{1 - k}$. Hence, the sequential rationality constraint of the agent:

$$(1 - \delta_{n-t+1})(k - 1) + \delta_{n-t+1}k = (1 - \delta_{n-t+1})k$$

Which is satisfied when $\delta_{n-t+1} = \frac{1}{1 + k}$.

In period $n - t$, $\frac{1 - 2k}{1 - k} \leq \rho_{n-t} < \frac{2 - 4k}{2 - 3k}$. Hence, the sequential rationality constraint of the agent:

$$(1 - \delta_{n-t})(k - 1) + \delta_{n-t}k = (1 - \delta_{n-t})k + \delta_{n-t} \frac{k^2 + \frac{1}{2}k}{1 + k}$$

Which is satisfied when $\delta_{n-t} = \frac{2 + 2k}{2 + 3k}$. Iterating in the same way, we get

$$\delta_i = \frac{2^{n-t+1-i}(1+2k) - 2k}{2^{n-t+1-i}(1+2k) - k}$$

and

$$\delta_1 = \frac{2^{n-t}(1+2k) - 2k}{2^{n-t}(1+2k) - k}$$

$\gamma_i = (1 - \delta_i) \prod_j^{i-1} \delta_j$ implies

$$\gamma_i = \frac{2^{i-1}k}{2^{n-t}(1+2k) - k}$$

Thus, in histories where no action 0 is observed, the stakes evolve as follows:

$$\Gamma = \left\{ \frac{k}{2^{n-t}(1+2k)-k}, \frac{2k}{2^{n-t}(1+2k)-k}, \dots, \frac{2^{n-t}k}{2^{n-t}(1+2k)-k} \right\}$$

□

Proof. **Proposition 2.3**

From Lemma 1.5, in any period i , if $\rho_{i+1}(0) < \frac{(1-2k)2^{n-i-2}}{(1-2k)2^{n-i-2}+k}$, then the good agent has higher reputational incentives, i.e. $\nu_i > 0$ implies $\mu_i = 1$. $\mu_i = 1$ can not be an equilibrium strategy as discussed in the proof of Proposition 2.2. Hence, such an equilibrium can exist only if $\nu_i = 0$ and $\mu_i \geq 0$. But then $\nu_i = 0$ implies $\rho_{i+1}(0) = 1$, contradiction. Then, in any equilibrium, $\rho_{i+1}(0) \geq \frac{(1-2k)2^{n-i-2}}{(1-2k)2^{n-i-2}+k}$ must hold.

Again following Lemma 1.5, if $\rho_{i+1}(0) \geq \frac{(1-2k)2^{n-i-2}}{(1-2k)2^{n-i-2}+k}$, then the bad agent has higher reputational concerns than the good agent. The reasoning is as follows. If $a_i = 0$ is observed, both type gets a future continuation payoff of $\delta_i k$. On the other hand, the expected continuation payoff of the good agent after $a_i = 1$ is higher than that of the bad agent. The reason is that in any period after $a_i = 1$, the good agent can increase her reputation costlessly with the probability $1/2$. Hence, $\delta_i^b < \delta_i^g$ where δ_i^b denotes the delta level for which the bad agent is indifferent between the two actions.

Following Proposition 2.2, the social planner's continuation payoff is increasing in δ_i if $\mu_i = 0 = \nu_i$. Thus, the payoff maximizing δ_i is whether δ_i^b or δ_i^g .

Case 1. $\delta_i = \delta_i^g$, $\mu_i \in [0, \mu'_i]$ and $\nu_i = 1$. Note that, from Lemma 1.1:

$$\nu_i < \frac{1 + \mu_i}{2} \leq \frac{1 + \frac{2k + \rho_i - 1}{\rho_i}}{2}$$

Also note that,

$$\frac{1 + \frac{2k + \rho_i - 1}{\rho_i}}{2} < 1$$

Hence, there is no equilibrium where $\mu_i \in [0, \mu'_i]$ and $\nu_i = 1$.

Case 2. $\delta_i = \delta_i^b, \mu_i = 0$ and $\nu_i \in [0, \frac{1}{2}]$ is the equilibrium strategy in the equilibrium. As we are focusing on the social planner optimal equilibrium, $\nu_i = 0$ for all i . It follows from the fact that the expected continuation payoff of the social planner is decreasing in ν_i .

Next, I will show by iteration that $\delta_i^b = \frac{1 + (n - t + 1 - i)k}{1 + (n - t + 2 - i)k}$.

In period $n - t + 1$, $1 - 2k \leq \rho_{n-t+1} < \frac{1 - 2k}{1 - k}$. Hence, the sequential rationality constraint of the bad agent:

$$(1 - \delta_{n-t+1})(k - 1) + \delta_{n-t+1}k = (1 - \delta_{n-t+1})k$$

Which is satisfied when $\delta_{n-t+1} = \frac{1}{1 + k}$.

In period $n - t$, $\frac{1 - 2k}{1 - k} \leq \rho_{n-t+1} < \frac{2 - 4k}{2 - 3k}$. Hence, the sequential rationality constraint of the bad agent:

$$(1 - \delta_{n-t})(k - 1) + \delta_{n-t}k = (1 - \delta_{n-t})k + \delta_{n-t}\frac{k^2}{1 + k}$$

Which is satisfied when $\delta_{n-t} = \frac{1 + k}{1 + 2k}$. Iterating in the same way, we get

$$\delta_i = \frac{1 + (n - t + 1 - i)k}{1 + (n - t + 2 - i)k}$$

and

$$\delta_1 = \frac{1 + (n - t)k}{1 + (n - t + 1)k}$$

$\gamma_i = (1 - \delta_i) \prod_j^{i-1} \delta_j$ implies

$$\gamma_i = \frac{k}{1 + (n - t + 1)k}$$

Thus, in histories where no action 0 is observed, each period receives the same stake. $\square$

Proof. **Proposition 2.4**

If the reputation in period i is sufficiently high - $\rho_i \geq \frac{(1-2k)2^{n-1-i}}{(1-2k)2^{n-1-i} + k}$ - then the principal's best response to $\mu_i = 0$ is $\lambda_i = 1$ and $\lambda_{i+1}(0) = \lambda_{i+1}(1) = 1$. In words, the principal will not interrupt the interaction after any action. Hence, the agent will set δ_i randomly in each period.

Consider (2).

Consider period i such that the agent has no reputation incentive. Note that, if $\mu_i = 0$, then the agent's expected continuation payoff is decreasing in δ_i , i.e. $\delta_i = \underline{\delta}$ is the optimal choice. In such i , the minimum expected continuation payoff the agent can get is given by $(1 - \underline{\delta})k$. On the other hand, if $\mu_i \geq 0$, the expected continuation payoff of the agent is increasing in δ_i , hence the choice of $\delta_i = \bar{\delta}$ in such i . Thus $\mu_i = 1$. I will show that some $\delta_i \geq \underline{\delta}$ and $\mu_i = 0$ leads to a higher expected continuation payoff for the agent compared to $\delta_i = \bar{\delta}$ and $\mu_i = 1$.

For $\delta_i = \bar{\delta}$ and $\mu_i = 1$, V_i^g is as following:

$$(1 - \bar{\delta})(k - 1) + \bar{\delta}k = k - \underline{\delta}$$

For $\frac{1}{1+k} > \delta_i \geq \underline{\delta}$ and $\mu_i = 0$, V_i^g is equal to or greater than the following:

$$(1 - \delta_i)k$$

Which is greater than $k - \underline{\delta}$ for $\delta_i k < \underline{\delta}$. Moreover, the agent should choose δ_i so that $V_i^P \geq 0$.

$$V_i^P = (1 - \delta_i) \left(k - \frac{1}{2} + \frac{1}{2}\rho_i \right) + \delta_i \frac{1}{2}\rho_i k \geq 0$$

$$\Leftrightarrow \delta_i \geq \frac{1 - 2k - \rho_i}{1 - 2k - \rho_i + \rho_i k}$$

$\delta_i^{min} = \frac{1 - 2k - \rho_i}{1 - 2k - \rho_i + \rho_i k}$ is the lowest δ_i for which the agent is hired. For lower delta levels, the future stakes are not high enough to compensate the negative payoff in the period i . As the agent's expected continuation payoff is decreasing in δ_i , the

choice of $\delta_i = \delta_i^{min} = \frac{1 - 2k - \rho_i}{1 - 2k - \rho_i + \rho_i k}$ as long as $\delta_i^{min} k \leq \underline{\delta}$. Note that, if $\delta_i^{min} \leq \underline{\delta}$, then $\delta_i = \underline{\delta}$ and $\underline{\delta} k \leq \underline{\delta}$ is trivially satisfied. If $\delta_i^{min} > \underline{\delta}$, then $\delta_i = \delta_i^{min}$ and the following constraint must hold:

$$\frac{1 - 2k - \rho_i}{1 - 2k - \rho_i + \rho_i k} k \leq \underline{\delta} \quad (\text{B.1})$$

Which holds if $\rho_i \geq \frac{(k - \underline{\delta})(1 - 2k)}{(k - \underline{\delta} + \underline{\delta} k)}$. Thus, for defined reputation levels, it is the best choice for the agent to choose δ_i^{min} and have no reputational incentives. But, if the reputation is lower than the defined level, then $\delta_i^{min} k > \underline{\delta}$. Then, it is optimal for the agent to set $\delta_i = \bar{\delta}$ and play action 0 with a probability 1. It is no surprise as the lower reputation implies the minimum delta the agent can set for the principal to hire is higher. If the minimum delta is high enough, then the agent rather prefers to set the highest possible stake to the future and screen her type.

Finally, consider (3).

I will define the lowest reputation level for which the agent is hired in each period. The expected continuation payoff of the principal is given by:

$$V_i^P = (1 - \bar{\delta}) \left(k - \frac{1}{2} \right) + \bar{\delta} \rho_i k \geq 0$$

$$\Leftrightarrow \rho_i \geq \frac{\underline{\delta}(1 - 2k)}{2(1 - \underline{\delta})k}$$

Hence, the agent is hired for the defined reputation levels.

□

Proof. **Proposition 2.5**

1. Follows from Proposition 1.1. As the agent is hired for every period, she is indifferent between different allocation of stakes. Hence, any $\delta_i \in [\underline{\delta}, \bar{\delta}]$ is an equilibrium strategy.

Now consider lower initial reputation levels.

Following Lemma defines two-periods equilibrium.

Lemma B.1. *If $n = 2$, following assessment constitutes the unique PBE. $\mu_1 = 0 =$*

ν_1

$$\delta_1 \in \begin{cases} [\underline{\delta}, \bar{\delta}], & \rho_1 \geq 1 - 2k \\ \delta'_1 = \frac{1 - 2k - \rho_1}{1 - 2k - \rho_1 + \rho_1 k}, & \frac{1}{2} - k \leq \rho_1 < 1 - 2k \end{cases}$$

$$\lambda_1 = \begin{cases} 0, & \rho_1 < \frac{1}{2} - k \\ 1, & \rho_1 \geq \frac{1}{2} - k \end{cases}$$

Proof. Lemma B.1

Again, $\rho_1 \geq 1 - 2k$ is obvious. Consider $\rho_1 < 1 - 2k$. Note that $\rho_2(1) < 1 - 2k$ in any agent-optimal equilibrium, hence $\lambda_2(1) = 0$. There are two possibilities:

Case 1

Agent has no reputation concern in the first period. In such case, V_1^g is decreasing in δ_1 . Hence, the agent chooses minimum possible δ_1 so that the principal gets a non-negative payoff.

$$V_1^P = (1 - \delta_1) \left(k - \frac{1}{2} + \frac{1}{2}\rho_1 \right) + \frac{1}{2}\rho_1 k \delta_1$$

V_1^P is non-negative if $\delta_1 \geq \delta'_1 := \frac{1 - 2k - \rho_1}{1 - 2k - \rho_1 + \rho_1 k}$. Thus, if there is no reputation concern in the lowest range of initial reputation so that the agent is hired, $\delta_1 = \delta'_1$ as the payoff of the agent is decreasing in δ_1 .

$\delta'_1 \leq \frac{1}{1+k}$ holds as long as $\rho_1 \geq \frac{1}{2} - k$. Thus, the agent is not hired for $\rho_1 < \frac{1}{2} - k$, if the agent has no reputation concern.

Case 2

Agent has reputation concern in the first period. Note that $\nu_1 = 1$ can not be an equilibrium strategy, because it implies $\rho_2(0) = \rho_1 < 1 - 2k$ by Bayes' rule and $\lambda_2(0) = 0$. Thus, $\nu_1 < 1$. Only δ_1 value satisfying this is $\frac{1}{1+k}$. Principal's payoff:

$$V_1^P = \frac{k}{1+k} \left(k - \frac{1}{2} + \frac{1}{2}\rho_1 \right) + \frac{1}{1+k} \frac{1}{2}\rho_1 k$$

Which is non-negative if $\rho_1 \geq \frac{1}{2} - k$. Hence, in any case, the agent is not hired if

$$\rho_1 < \frac{1}{2} - k.$$

Note that:

$$V_1^{g[case1]} = (1 - \delta'_1)k + \frac{1}{2}\delta'_1 k$$

$$V_1^{g[case2]} = \frac{k^2 + \frac{1}{2}k}{1 + k}$$

and $\delta'_1 \leq \frac{1}{1+k}$ implies $V_1^{g[case1]} \geq V_1^{g[case2]}$, which holds strictly if $\delta'_1 < \frac{1}{1+k}$. Thus, no reputation concern setup is preferred by the agent. $\square$

I will prove the rest with induction method. Consider period $n - 2$. If, $\rho_{n-2} \geq 1 - 2k$, $\delta_{n-2} = \underline{\delta}$. Consider $\rho_{n-2} < 1 - 2k$. From Lemma B.1, $V_{n-1}^P = 0$ if $\rho_{n-1} < 1 - 2k$. As $\rho_{n-2} < 1 - 2k$ implies $\rho_{n-1}(1) < 1 - 2k$, $V_{n-1}^P(|a_n - 2| = 1) = 0$.

Consider two cases:

Case 1

Agents sets δ_{n-2} so that there is no reputation concern.

$$V_{n-2}^P = (1 - \delta_{n-2}) \left(k - \frac{1}{2} + \frac{1}{2}\rho_{n-2} \right) + \delta_{n-2} \frac{1}{2}\rho_{n-2}k$$

Thus, the agent sets minimum δ_{n-2} providing non-negative payoff to the principals,

i.e. $\delta_{n-2} = \delta'_{n-2}$. Agent is not hired if $\rho_{n-2} < \frac{1}{2} - k$.

Case 2

Agent sets δ_{n-2} so that there are reputation concerns.

Case 2.1

$\rho_{n-1}(0) \geq \frac{1 - 2k}{1 - k}$. This implies that if $a_i = 0$ is observed in period $n - 2$, the agent will be hired for the rest of the game irrespective of a_{n-1} . Define $\tilde{\delta}_{n-2}$ which makes the agent indifferent. If $\mu_{n-2} = 1$, then $\rho_{n-1}(1) = 0$. In this case, $\lambda_{n-1}(1) = 0$

and $\tilde{\delta}_{n-2} = \frac{1}{1+k}$. Thus, the problem is the same with two periods interaction. I showed in Lemma B.1 that the agent prefers no reputation setup. Hence, $\mu_{n-2} < 1$. Good agent is either indifferent or strictly prefers action 1. Hence, for the sake of exposition, assume $\mu_{n-2} = 0$, wlog.

$$V_{n-2}^{g[rep]} = (1 - \tilde{\delta}_{n-2})k + \frac{1}{2}\tilde{\delta}_{n-2}k + \frac{1}{2}\tilde{\delta}_{n-2}V_{n-1}^{g[rep]}(|a_{n-2} = 1)$$

where $V_{n-1}^{g[rep]}(|a_{n-2} = 1) = 0$ if $\rho_{n-1} < \frac{1}{2} - k$ and

$$(1 - \delta'_{n-1})k + \frac{1}{2}\delta'_{n-1}k$$

otherwise.

On the other hand, if the agent sets $\delta_{n-2} = \delta'_{n-2}$ as in case 1, she gets:

$$V_{n-2}^{g[norep]} = (1 - \delta'_{n-2})k + \frac{1}{2}\delta'_{n-2}k + \frac{1}{2}\delta'_{n-2}V_{n-1}^{g[norep]}(|a_{n-2} = 1)$$

Note that, $\delta'_{n-2} \leq \tilde{\delta}_{n-2}$ and $V_{n-1}^{g[norep]}(|a_{n-2} = 1) \geq V_{n-1}^{g[rep]}(|a_{n-2} = 1)$ (as $\rho_{n-1}(1)$ decreases with μ_{n-2}). Hence, agent prefers to set δ'_{n-2} .

Case 2.2

Hence, if there are reputation concerns in the equilibrium in period $n-2$, it must be so that $1 - 2k \leq \rho_{n-1}(0) < \frac{1}{2} - k$. We need $\nu_{n-2} \in (0, 1)$ by Bayes' rule. Then $\mu_{n-2} = 1$. Note that, if $\mu_{n-2} = 1$, $\rho_{n-1}(1) = 0$ and $\lambda_{n-1}(1) = 0$. Hence, δ_{n-2} is such that the bad agent is indifferent:

$$(1 - \delta_{n-2})k = (1 - \delta_{n-2})(k - 1) + \delta_{n-2}(1 - \underline{\delta})k$$

which implies:

$$\delta_{n-2} = \frac{1}{1 + k - \underline{\delta}}$$

The good agent's payoff is given by:

$$V_{n-2}^{g[case2]} = \frac{k - \underline{\delta}k}{1 + k - \underline{\delta}k} \left(k - \frac{1}{2} \right) + \frac{1}{1 + k - \underline{\delta}k} \left(k - \frac{1}{2}\underline{\delta}k \right)$$

The good agent's payoff in case 1 where there are no reputation concerns:

$$V_{n-2}^{g[case1]} = (1 - \delta'_{n-2})k + \frac{1}{2}\delta'_{n-2}k + \frac{1}{2}\delta'_{n-2}V_{n-1}^g(|a_{n-2} = 1)$$

Note that $V_{n-1}^g(|a_{n-2} = 1) < k$, hence $V_{n-2}^{g[case1]}$ is decreasing in δ'_{n-2} . It implies that

$V_{n-2}^{g[case1]} > \frac{k^2 + \frac{1}{2}k}{1+k} > V_{n-2}^{g[case2]}$. Hence the agent prefers to set δ_{n-2} to the lowest possible value and there are no reputation concerns.

Finally, consider period i such that strategies are as defined above in periods $i+1, \dots, n$. If, $\rho_i \geq 1 - 2k$, $\delta_i = \underline{\delta}$. Consider $\rho_i < 1 - 2k$. From above analysis, $V_{i+1}^P = 0$ if $\rho_{i+1} < 1 - 2k$. $\rho_i < 1 - 2k$ implies $\rho_{i+1}(1) < 1 - 2k$, $V_{i+1}^P(1) = 0$. So, following the same steps above, if agent sets δ_i so that there are no reputation concerns, then $\delta_i = \delta'_i$.

Now, suppose the agent sets δ_i so that there are reputation concerns. Define $s, t \in \{1, 2, \dots, n-i\}$ such that $s < t$. Let t represent the number of periods the agent is hired if action 1 is observed in all t periods following period $a_i = 0$. Following the reasoning above, the agent will not set $t = n-i$. Hence, if there are reputation concerns, then $t < n-i$. It implies good agent has higher reputation concerns, hence $\mu_i = 1$ and $\nu_i \in (0, 1)$. Note that, t represents the number of periods after which the reputation falls below $\frac{1}{2} - k$ if no action 0 is observed. Let s represent the number of periods after which the reputation falls below $1 - 2k$. The reason for making this restriction is that in any period the agent's reputation fall below $1 - 2k$, the agent can no more set $\delta_j = \underline{\delta}$ but rather $\delta_j = \delta'_j$.

Define Z as the payoff of the agent in these t periods. Note that, both types of agents get a payoff of k in each of these periods. Then Z has the following form:

$$Z := (1 - \underline{\delta}^{t-s})k + \underline{\delta}^{t-s}\tilde{V}$$

Where $\tilde{V}$ stands for the payoff in the last s periods. $\tilde{V} < k$ as agent sets $\delta_j = \delta'_j$ in these periods. Hence, $Z \leq k$. In order to find the form of agent's payoff, we first

need to find δ_i . Note that, $\nu_i \in (0, 1)$, hence, the sequential rationality constraint of the bad agent:

$$(1 - \delta_i)k = (1 - \delta_i)(k - 1) + \delta_i Z$$

Hence, $\delta_i = \frac{1}{1 + Z}$.

Payoff of the good agent becomes:

$$V_i^{g[case2]} = \frac{Z}{1 + Z} \left(k - \frac{1}{2} \right) + \frac{1}{2^t} \frac{Z}{1 + Z} + \left(1 - \frac{1}{2^t} \right) \frac{k}{1 + Z}$$

Where $\frac{1}{2^t}$ is the probability that no action 0 is observed in t periods. With the complementary probability, the agent reveals herself as the good type and gets payoff of k . $V_i^{g[case2]}$ is continuous and monotonically increases in $\underline{\delta}$. Moreover, it is equal to $\frac{k^2 + 1/2k}{1 + k}$ if $\underline{\delta} = 0$. Recall that $V_{n-2}^{g[case1]} > \frac{k^2 + 1/2k}{1 + k}$, hence, if $\underline{\delta}$ is sufficiently small, then the good agent prefers to set stakes so that there are no reputation concerns. $\square$

Proof. Proposition 2.6

Suppose there is an equilibrium where $\frac{2^{t+1}(1 - 2k)}{2^{t+1}(1 - 2k) + k} \leq \rho_1$ and the principal sets $\delta_1 < \frac{1}{1 + k}$.

$$V_1^P = (1 - \delta_1) \left(k - \frac{1}{2} + \frac{1}{2}\rho_1 \right) + \delta_1 V_2^P$$

By Bayes' rule, the agent is hired for at least $t + 2$ periods in the equilibrium. In any following period, the principal can set δ values so that the agent has no reputation concern in any period. Thus,

$$V_2^P \geq \frac{1}{2}\rho_1 k + \frac{1}{2} \left(1 - \frac{1}{2^{t+1}} \right) \rho_1 k$$

Which is greater than $k - \frac{1}{2} + \frac{1}{2}\rho_1$, by $\rho_1 < \frac{2^{t+1}(1-2k)}{2^{t+1}(1-2k) + k}$. Hence, V_1^P is increasing in δ_1 . Thus, the principal will deviate to a higher δ_1 . Hence, in any equilibrium, $\delta_1 \geq \frac{1}{1+k}$. The same reasoning applies to period i . $\square$

Proof. **Proposition 2.7**

I will use the induction method for proof. Consider period i where ρ_i is sufficiently low, i.e. $\lambda_{i+1}(1) = 0$. Suppose the principal sets δ_i such that $\mu_i = 0$. The expected continuation payoff of the principal:

$$(1 - \delta_i) \left(k - \frac{1}{2} + \frac{1}{2}\rho_i \right) + \delta_i \frac{1}{2}\rho_i k$$

Which is increasing in δ_i . From the sequential rationality constraint of the agent,

$\delta_i \leq \frac{1}{1+k}$. Hence, $\delta_i = \frac{1}{1+k}$ maximizes the expected payoff of the principal if $\mu_i = 0$. Note that, the principal's expected payoff does not depend on μ_i if $\delta_i = \frac{1}{1+k}$. Hence, the maximum expected payoff of the principal if $\mu_i \in [0, 1]$ is given by

$$\frac{k}{1+k} \left(k - \frac{1}{2} + \rho_i \right) \tag{B.2}$$

On the other hand, if $\mu_i = 1$, then the principal's expected payoff is increasing in δ_i . The payoff maximizing choice of delta: $\delta_i = \bar{\delta}$. The payoff becomes:

$$(1 - \bar{\delta}) \left(k - \frac{1}{2} \right) + \bar{\delta}\rho_i k \tag{B.3}$$

For $\bar{\delta}$ sufficiently close to 1, V_i^P is sufficiently close to $\rho_i k$. (B.3) > (B.2) and $\delta_i = \bar{\delta}$ leads to a higher expected payoff to the principal. $V_i^g = (1 - \bar{\delta}) \left(k - \frac{1}{2} \right) + \bar{\delta}k =$

$$k - \frac{1}{2} + \frac{1}{2}\bar{\delta}.$$

Consider period $i - 1$. As $\rho_i(0) = 1$, the sequential rationality constraint of the good agent for $\mu_{i-1} \geq 0$:

$$(1 - \delta_{i-1})(k - 1) + \delta_{i-1}k \geq (1 - \delta_{i-1})k + \delta_{i-1} \left(k - \frac{1}{2} + \frac{1}{2}\bar{\delta} \right)$$

$$\delta_{i-1} \geq \frac{2}{3 - \bar{\delta}}$$

Which can not hold for $\delta_{i-1} < \bar{\delta}$. The good agent has reputational incentive only if she is not hired after action 1, i.e. $\mu_{i-1} = 1$. Then the principal sets $\delta_{i-1} = \bar{\delta}$. The principal gets the following expected continuation payoff:

$$(1 - \bar{\delta}) \left(k - \frac{1}{2} \right) + \bar{\delta} \rho_{i-1} k \quad (\text{B.4})$$

On the other hand, following the sequential rationality constraint of the agent, $\mu_{i-1} = 0$ is also a best response to $\delta_{i-1} = \bar{\delta}$ which leads to the expected continuation payoff for the principal approximated as such:

$$\begin{aligned} & (1 - \bar{\delta}) \left(k - \frac{1}{2} + \rho_{i-1} \right) + \bar{\delta} \frac{1}{2} \rho_{i-1} k + \bar{\delta} \left(1 - \frac{1}{2} \rho_{i-1} \right) \rho_i(1) k \\ & (1 - \bar{\delta}) \left(k - \frac{1}{2} + \rho_{i-1} \right) + \bar{\delta} \rho_{i-1} k \end{aligned} \quad (\text{B.5})$$

The payoff in (B.5) is greater than (B.4), hence the $\mu_{i-1} = 0$ and $\delta_{i-1} = \bar{\delta}$ in the equilibrium.

Consider any period $j < i - 1$ such $\delta_{j+1} = \bar{\delta}$ and $\mu_{j+1} = 0$. Following the reasoning above, the agent has reputational incentive only if she is not hired after action 1. Hence $\mu_j = 1$ and $\delta_j = \bar{\delta}$ and the principal gets

$$(1 - \bar{\delta}) \left(k - \frac{1}{2} \right) + \bar{\delta} \rho_{i-1} k \quad (\text{B.6})$$

On the other hand, following the sequential rationality constraint of the agent, $\mu_j = 0$ is also a best response to $\delta_j = \bar{\delta}$ which leads to the expected continuation payoff for the principal approximated as such:

$$(1 - \bar{\delta}) \left(k - \frac{1}{2} + \rho_j \right) + \bar{\delta} \rho_j k \quad (\text{B.7})$$

Which is greater than the payoff in B.6. Hence, $\mu_j = 0$.

Thus, in the equilibrium, the principal sets $\delta_i = \bar{\delta}$ for all $i \in N$. $\mu_i = 0$ for all i except when the reputation is sufficiently small so that $\lambda_{i+1}(1) = 0$ where $\mu_i = 1$. $\square$

Proof. Proposition 2.8

The following Lemma will be used in the proof.

Lemma B.2. *If $\rho_1 \leq \frac{(1 - 2k)2^{n-2}}{(1 - 2k)2^{n-2} + k}$, then there is an evaluation phase in the agent-optimal equilibrium i.e. $\lambda_i(0) = 1$ and $\lambda_i(1) = 0$ for some $i \in N$, and the strategy of the agent in the evaluation phase is as following: $\mu_i = 1$ and $\nu_i \in [0, 1]$. The last period of the evaluation phase may be an exception where $\mu_i \in [0, 1]$ and $\nu_i = 0$ may hold.*

Proof. Lemma B.2.

Suppose that there is an equilibrium where there is no evaluation phase. If $\lambda_i(0) = 1$ and $\lambda_i(1) = 0$ for all $i \in N$ in the equilibrium, then the agent has no incentive to play action 0 with a positive probability: $\mu_i = 0$ and $\nu_i = 0$ for all $i \in N$. By Bayes' rule, $\rho_1 \leq \frac{(1 - 2k)2^{n-2}}{(1 - 2k)2^{n-2} + k}$ implies $\rho_n(1) < 1 - 2k$ if $a_i = 1$ is observed in each period $i \in N/n$. Thus, $\lambda_n(1) = 0$, contradiction.

Next, I will show that $\mu_i = 1$ and $\nu_i \in [0, 1]$ in the evaluation phase. In any period where action 0 does not lead to the agent being hired for all the remaining periods, then the good agent has strictly higher reputational concerns than the bad agent. Fix a period i so that $\lambda_{i+1}(0) = 1$ and $\lambda_{i+1}(1) = 0$. Then the expected continuation payoff of the good agent is higher than that of bad agent in any such

period i . Consider the following: the bad agent can improve her reputation by playing the action 0, but it costs her today's payoff. So she sacrifices that period's gain in order to get more chance for further periods. On the other hand, suppose $\theta_i = 1$ is observed and the good agent plays $a_i = 0$ in period i . She sacrifices period i payoff but in return she gets more chance to play her favorite action in the future. In addition, the good agent gets more chance to improve her reputation costlessly in the future (when the true state is 0 and she plays accordingly). Consider the following. Suppose the evaluation phase also involves the period $i + 1$. In period $i + 1$, the expected payoff of the good agent from playing action 0 is given by $k - \text{prob}[\theta_{i+1} = 1]1 - \text{prob}[\theta_{i+1} = 0]0 = k - \frac{1}{2}$. Whereas, the expected payoff of the bad type from playing action 0 in period $i + 1$ is given by $k - \text{prob}[\theta_{i+1} = 1]1 - \text{prob}[\theta_{i+1} = 0]1 = k - 1$. Thus, building reputation is costlier for the bad agent in period i . If period i is not the last period of the evaluation phase, the good agent's expected continuation payoff is always higher than that of the bad agent. If period i is the last period of the evaluation phase so that there is a perfect screening in the period, then the future expected payoff is the same for both types.

Suppose $\nu_i = 0$, then $\rho_{i+1}(0) = 1$. Thus, if $\nu_i = 0$, then the period i serves as the evaluation period, and the agent is hired until the end if action 0 is observed in period i . It is possible in the last period of the evaluation phase and $\mu_i \in [0, 1]$ satisfies the equilibrium condition. The reasoning is that the last period of the evaluation phase, the decision is exactly same with the two-periods game. If i is not the last period of the evaluation phase, i.e. $\lambda_{i+1}(0) = 1$ and $\lambda_{i+1}(1) = 0$, then $\nu_i > 0$.

Suppose $\nu_i = 1$, it implies $\mu_i = 1$. Then there will be no update on the type of the agent and $\rho_{i+1}(0) = \rho_i$. The reputation is not updated and the principal gets a negative expected stage game payoff in period i . As $\delta_i < 1$, $\mu_i = 1 = \nu_i$ can not be a part of the agent-optimal equilibrium. Hence, we need to have $\nu_i \in (0, 1)$ in the evaluation phase may be except from the last period. From the reasoning above, the good agent's future expected continuation payoff is larger than that of the bad agent. For the δ_i level that the bad agent is indifferent, the good agent strictly

prefers playing action 0.

Hence, if $\rho_1 \leq \frac{(1-2k)2^{n-2}}{(1-2k)2^{n-2}+k}$, then there will be evaluation phase in any equilibria and the equilibrium strategy of the agent in the evaluation phase is given by $\mu_i = 1$ and $\nu_i \in (0, 1)$ maybe expect from the last period. $\square$

Consider the following strategies:

$$\nu_i \in \begin{cases} \{0\}, & \delta_i < \frac{1}{1+k} \\ [0, 1], & \frac{1}{1+k} \leq \delta_i \leq \frac{\sum_{h=0}^{n-i-1} k^h}{\sum_{h=0}^{n-i} k^h} \\ \{1\}, & \delta_i > \frac{\sum_{h=0}^{n-i-1} k^h}{\sum_{h=0}^{n-i} k^h} \end{cases}$$

$$\mu_i \in \begin{cases} \{0\}, & \delta_i < \frac{1}{1+k} \\ [0, 1], & \frac{1}{1+k} \leq \delta_i \leq \delta_i^g \\ \{1\} & \delta_i > \delta_i^g \end{cases}$$

$$\text{such that } \frac{1}{1+k} \leq \delta_i^g \leq \frac{\sum_{h=0}^{n-i-1} k^h}{\sum_{h=0}^{n-i} k^h} \text{ for } i \in N/n, \mu_n = 0, \nu_n = 0.$$

Suppose that the initial reputation level ρ_1 is close to 0. There will be an evaluation phase from Lemma B.2. If ρ_1 is the smallest initial reputation so that the project is run, then the evaluation phase will start in the first period. The reasoning is that, if $\lambda_2(0) = 1 = \lambda_2(1)$, then $\rho_2(1) < \rho_1$. As ρ_1 is the smallest reputation level for the project to be implemented, $\rho_2(1) < \rho_1$ implies $\lambda_2(1) = 0$, a contradiction. Suppose that the first $m-1$ periods serve as the evaluation phase. Then $\lambda_m(0) = 1$ and $\lambda_m(1) = 0$ and $\lambda_{m+1}(0) = 1$ and $\lambda_{m+1}(1) = 1$. In the evaluation phase, the strategy of the agent: $\mu_i = 1$ and $\nu_i \in (0, 1)$ may be except from the last period. Thus, in each period of the evaluation phase the principal gets an expected payoff of

$k - \frac{1}{2}$. Define $q_i := \text{prob}[a_i = 0] = \rho_i (1 + \mu_i) + (1 - \rho_i) \nu_i$. The expected continuation payoff of the principal is given as follows:

$$V_1^P = \gamma_1 \left(k - \frac{1}{2} \right) + \gamma_2 q_1 \left(k - \frac{1}{2} \right) + \gamma_3 q_1 q_2 \left(k - \frac{1}{2} \right) + \dots$$

$$\dots + \gamma_{m-1} \prod_{i=1}^{m-2} q_i \left(k - \frac{1}{2} \right) + \delta_m \prod_{i=1}^{m-1} q_i V_m^P$$

Note that V_1^P is increasing in ρ_1 and $V_1^P = \gamma_1 \left(k - \frac{1}{2} \right)$ if $\rho_1 = 0$.

If the period $m - 1$ is the last period of the evaluation phase, then the sequential rationality constraint of the bad agent in period $m - 1$ is given as:

$$(1 - \delta_{m-1}) k = \delta_{m-1} (k - 1) + (1 - \delta_{m-1}) f(k)$$

Where $f(k) = V_m^B$ is increasing function of k such that $f(k) \leq k$. The sequential

rationality constraint implies $\delta_{m-1} = \frac{1}{1 + f(k)}$. The sequential rationality constraint of the bad agent in period $m - 2$ is given as:

$$(1 - \delta_{m-2}) k = \delta_{m-2} (k - 1) + (1 - \delta_{m-2}) \frac{f(k) k}{1 + f(k)}$$

Hence,

$$\delta_{m-2} = \frac{1 + f(k)}{1 + f(k) + f(k) k}$$

Iterating in the same fashion,

$$\delta_i = \frac{1 + \sum_{h=1}^{m-i-1} f(k) k^{h-1}}{1 + \sum_{h=1}^{m-i} f(k) k^{h-1}} \text{ for all } i \in \{0, 1, \dots, m-2\}$$

and

$$\delta_{m-1} = \frac{1}{1 + f(k)}$$

Thus,

$$\delta_1 = \frac{1 + \sum_{h=1}^{m-2} f(k)k^{h-1}}{1 + \sum_{h=1}^{m-1} f(k)k^{h-1}}$$

By definition,

$$\gamma_1 = 1 - \delta_1 = \frac{f(k)k^{m-2}}{1 + \sum_{h=1}^{m-1} f(k)k^{h-1}}$$

Given $f(k) \leq k < \frac{1}{2}$, γ_1 is decreasing in m . $\left(k - \frac{1}{2}\right)$ is decreasing in γ_1 and is maximized when $m = n$. Increasing the span of the evaluation phase increases the expected continuation payoff of the principal if the initial reputation is close enough to 0. Hence, in the principal optimal equilibrium, the evaluation phase includes all the periods except period n .

Note that, if $m = n$, then $f(k) = k$, hence

$$\delta_i^* = \frac{1 + \sum_{h=1}^{n-i-1} f(k)k^{h-1}}{1 + \sum_{h=1}^{n-i} f(k)k^{h-1}} \text{ and } \lambda_i^*(0) = 1 \text{ and } \lambda_i^*(1) = 0 \text{ for all } i \in N$$

Note that $\gamma_i = (1 - \delta_i) \prod_j^{i-1} \delta_j$ thus,

$$\gamma_i = \frac{1 - k}{1 - k^n} k^{n-i}$$

The last thing to show is that there exist small enough initial reputation level that provides a positive expected continuation payoff to the principal for the above-defined strategies.

The expected continuation payoff of the principal is given below:

$$V^{P*} = \gamma_1 \left(k - \frac{1}{2} \right) + \gamma_2 q_1 \left(k - \frac{1}{2} \right) + \gamma_3 q_1 q_2 \left(k - \frac{1}{2} \right) + \dots$$

$$\dots + \gamma_{n-1} \prod_{i=1}^{n-2} q_i \left(k - \frac{1}{2} \right) + \gamma_n \prod_{i=1}^{n-1} q_i \left(k - \frac{1}{2} + \frac{1}{2} \rho_n \right)$$

Note that the highest ρ_{n-1} satisfying the equilibrium conditions is $\frac{1-2k}{1-k}$. Hence,

in the principal optimal equilibrium, $\nu_{n-2} = \frac{k\rho_{n-2}}{(1-2k)(1-\rho_{n-2})}$. I also know from the two-periods analysis that $\nu_{n-1} = 0$. For other ν_i values, the principal's utility is decreasing in ν_i . Thus, the expected continuation payoff of the principal will be higher than the case where $\nu_i = 1$ for all $i \in N/\{n-2, n-1, n\}$. Define the expected continuation payoff of the principal as $\underline{V}^{P*}$ as the lower bound of the payoff he will get, which holds when $\nu_i = 1$ for all $i \in N/\{n-2, n-1, n\}$.

Simplify the payoff:

$$\underline{V}^{P*} = \frac{1-k}{1-k^n} \left(k^{n-1} \left(k - \frac{1}{2} \right) + k^{n-2} \left(k - \frac{1}{2} \right) + k^{n-3} \left(k - \frac{1}{2} \right) + \dots \right.$$

$$\left. \dots + k^2 \left(k - \frac{1}{2} \right) + k \left(\rho_1 + (1-\rho_1) \frac{k\rho_1}{(1-2k)(1-\rho_1)} \right) \left(k - \frac{1}{2} \right) + \rho_1 k \right)$$

where I used the fact that, if $\nu_i = 1$ for all $i \in N/\{n-2, n-1, n\}$, there is no update on the beliefs until the last two periods. Thus, $\text{prob}[a_{n-2} = 0] =$

$\rho_{n-2} + (1-\rho_{n-2})\nu_{n-2} = \rho_1 + (1-\rho_1) \frac{k\rho_1}{(1-2k)(1-\rho_1)}$. Moreover, as the principal's payoff does not depend on μ_{n-1} , I assume $\mu_{n-1} = 1$ for the sake of simplicity. Then,

$\text{prob}[a_{n-1} = 0] = \rho_{n-1} + (1-\rho_{n-1})\nu_{n-1} = \frac{\rho_1}{\rho_1 + (1-\rho_1)\nu_{n-2}}$ and $\rho_n = 1$ as $\nu_{n-1} = 0$.

The payoff becomes:

$$\underline{V}^{P*} = \frac{1-k}{1-k^n} \left\{ \frac{k^2 - k^n}{1-k} \left(k - \frac{1}{2} \right) + \frac{1}{2} \rho_1 k (1+k) \right\}$$

which is non-negative for: $\rho_1 \geq \frac{(1-2k)(k-k^{n-1})}{(1-k^2)}$. Hence there exist such ρ_1 where V^{P*} is non-negative. $\square$

Appendix C

CHAPTER 3 COMPARISON OF TREATMENT PARAMETERS

Given a treatment, the two payoff tables (one for the principal and the good agent, and the other for the bad agent) that are employed for each period (stage) are given below. The first-period payoff tables are the same in all treatments. The second- and third-period payoff tables are adjusted for the defined stage stakes in each treatment.

Table C.1: Treatment 1 payoff table.

Principal / Type 1 Agent			Type 2 Agent		
	$\theta = 0$	$\theta = 1$		$\theta = 0$	$\theta = 1$
$a = 0$	10	-20	$a = 0$	-20	-20
$a = 1$	-20	10	$a = 1$	10	10

Treatment 1 is the baseline treatment at which the weights (stakes) in all periods are equal, i.e., $\Gamma_1 = \{\frac{1}{3}, \frac{1}{3}, \frac{1}{3}\}$. That is, the payoff table is employed in each period in the game.

Period 1

Principal / Type 1 Agent			Type 2 Agent		
	$\theta = 0$	$\theta = 1$		$\theta = 0$	$\theta = 1$
$a = 0$	10	-20	$a = 0$	-20	-20
$a = 1$	-20	10	$a = 1$	10	10

Period 2

Principal / Type 1 Agent			Type 2 Agent		
	$\theta = 0$	$\theta = 1$		$\theta = 0$	$\theta = 1$
$a = 0$	20	-40	$a = 0$	-40	-40
$a = 1$	-40	20	$a = 1$	20	20

Period 3

Principal / Type 1 Agent			Type 2 Agent		
	$\theta = 0$	$\theta = 1$		$\theta = 0$	$\theta = 1$
$a = 0$	40	-80	$a = 0$	-80	-80
$a = 1$	-80	40	$a = 1$	40	40

$$\Gamma_2 = \left\{ \frac{1}{7}, \frac{2}{7}, \frac{4}{7} \right\}.$$

Table C.2: Treatment 2 payoff table.

Table C.3: Treatment 3 payoff table.

Period 1

Principal / Type 1 Agent			Type 2 Agent		
	$\theta = 0$	$\theta = 1$		$\theta = 0$	$\theta = 1$
$a = 0$	10	-20	$a = 0$	-20	-20
$a = 1$	-20	10	$a = 1$	10	10

Period 2

Principal / Type 1 Agent			Type 2 Agent		
	$\theta = 0$	$\theta = 1$		$\theta = 0$	$\theta = 1$
$a = 0$	30	-60	$a = 0$	-60	-60
$a = 1$	-60	30	$a = 1$	30	30

Period 3

Principal / Type 1 Agent			Type 2 Agent		
	$\theta = 0$	$\theta = 1$		$\theta = 0$	$\theta = 1$
$a = 0$	90	-180	$a = 0$	-180	-180
$a = 1$	-180	90	$a = 1$	90	90

$$\Gamma_3 = \left\{ \frac{1}{13}, \frac{3}{13}, \frac{9}{13} \right\}.$$

Table C.4: Treatment 4 payoff table.

Period 1

Principal / Type 1 Agent			Type 2 Agent		
	$\theta = 0$	$\theta = 1$		$\theta = 0$	$\theta = 1$
$a = 0$	10	-20	$a = 0$	-20	-20
$a = 1$	-20	10	$a = 1$	10	10

Period 2

Principal / Type 1 Agent			Type 2 Agent		
	$\theta = 0$	$\theta = 1$		$\theta = 0$	$\theta = 1$
$a = 0$	40	-80	$a = 0$	-80	-80
$a = 1$	-80	40	$a = 1$	40	40

Period 3

Principal / Type 1 Agent			Type 2 Agent		
	$\theta = 0$	$\theta = 1$		$\theta = 0$	$\theta = 1$
$a = 0$	160	-320	$a = 0$	-320	-320
$a = 1$	-320	160	$a = 1$	160	160

$$\Gamma_4 = \left\{ \frac{1}{21}, \frac{4}{21}, \frac{16}{21} \right\}.$$

Table C.5: Treatment parameters.

	Treatment 1	Treatment 2	Treatment 3	Treatment 4
Stakes	$\left\{ \begin{matrix} 1 & 1 & 1 \\ \bar{3}' & \bar{3}' & \bar{3}' \end{matrix} \right\}$	$\left\{ \begin{matrix} 1 & 2 & 4 \\ \bar{7}' & \bar{7}' & \bar{7}' \end{matrix} \right\}$	$\left\{ \begin{matrix} 1 & 3 & 9 \\ \bar{13}' & \bar{13}' & \bar{13}' \end{matrix} \right\}$	$\left\{ \begin{matrix} 1 & 4 & 16 \\ \bar{21}' & \bar{21}' & \bar{21}' \end{matrix} \right\}$
Show-up fee	20TL	20TL	20TL	20TL
Exchange rate - e	1/3	1/7	1/13	1/21
Min. ECU/super-game	-60	-140	-260	-420
Min. TL/super-game	-20	-20	-20	-20
Max. ECU/super-game	30	70	130	210
Max. TL/super-game	10	10	10	10
Quiz ECU/correct answer	6	14	26	42
Quiz TL/correct answer	2	2	2	2
Survey payment	10TL	10TL	10TL	10TL
Halt-Laury task (ECU)	A: 36 vs 27	A: 90 vs 60	A: 160 vs 120	A: 250 vs 190
	B: 60 vs 9	B: 140 vs 20	B: 260 vs 40	B: 420 vs 60
Halt-Laury task (TL)	A: 12 vs 9	A: 12.86 vs 8.6	A: 12.31 vs 9.23	A: 11.9 vs 9.05
	B: 20 vs 3	B: 20 vs 2.86	B: 20 vs 3.07	B: 20 vs 2.86

Appendix D

CHAPTER 3 INSTRUCTIONS AND SCREENSHOTS (TREATMENT 4)

D.1 English Instructions

INTRO

This is a decision-making experiment and part of a scientific project.

This experiment aims to understand how people make decisions in different situations. Your decisions will not be evaluated as "right" or "wrong".

The payment you will get in the experiment depends on your decisions, and how your earnings will be determined is detailed in this instruction. Therefore, you must read and understand the instruction carefully. The payment of your earnings will be made in the BELIS laboratory in cash or in the form of a money transfer to your bank account. The other participants will not be informed about your earnings.

The decisions you make and the answers you give are entirely anonymous, and they are not matched with any identity information. You are strictly prohibited from communicating with other participants until the experiment ends.

(Checker) I confirm that I will not communicate with other participants during the experiment.

During the experiment, your communication with the experiment administrators will be via the Zoom session, the link of which has been shared with you. Experiment administrators will make audio announcements and answer any questions you may have through this session. Therefore, it is essential to stay connected to Zoom and keep your device's volume unmuted to hear the announcements during the experiment.

You can forward your questions to the experiment administrators using the

“chat” feature of the Zoom session.

(Checker) I will not leave the Zoom session during the experiment and keep the device volume up to hear the announcements.

The experiment takes approximately 120 min.

Your experiment links have been created individually. Please do not share your link with anyone.

(Checker) I confirm that I will not share my test link with anyone.

If you experience a disconnection during the experiment, you can log in via the same link and continue where you left off. If your problem persists, please contact one of the experimenters via Zoom.

You must not leave the experiment application and Zoom session before the end of the experiment. Even if a single participant does not complete the experiment, the experiment data will become unusable for the project team.

If you leave the experiment before the end of the experiment, no payment will be made to you, and you will not be able to participate in the future BELIS experiments.

(Checker) I confirm that I will not leave the experiment until the experiment is completed.

General Information and Roles.

The computer will randomly select one of the Green and Blue colors in each Round. An employer who cannot observe the color chosen by the computer decides whether or not to give the job of observing this choice on his behalf and making a decision based on it to the employee with whom he is matched. If he gives the job, he decides whether to continue working with the employee based on his decision.

The employee is either Type 1 or Type 2.

- If Type 1, he will benefit both himself and the employer if he chooses the same color as the computer.
- If Type 2: regardless of the computer’s choice, he gets a profit if he chooses the color blue.

The employer does not know whether the employee is Type 1 or Type 2.

You will be randomly assigned to either the employer or employee roles when the experiment starts. If you are assigned to the employee role, you will be Type 1 with 10 % probability and Type 2 with 90 % probability.

Your assigned role will not change until the end of the experiment. If you were appointed as the employer at the beginning of the experiment, you would remain the employer until the experiment ends. If you were appointed as the Type 1 employee at the beginning of the experiment, you would remain the Type 1 employee until the experiment ends. If you were appointed as the Type 2 employee at the beginning of the experiment, you would remain the Type 2 employee until the experiment ends.

Sections and Rounds.

An employer and an employee are matched at the start of each section. The employer does not know whether the employee he is matched with is Type 1 or Type 2. He knows that the probability of the employee being Type 1 is 10%, and the probability of being Type 2 is 90%.

Each section contains a maximum of 3 rounds. At the beginning of each round, the employer decides whether to assign the job to the matched employee. If the employer decides not to give the job, the section ends regardless of the Round number. For the next section, the employer is randomly matched with an employee.

How will your section earnings be determined?

The earnings the employer and the employee earned in the rounds completed within a section determine their section earnings.

Example 1: Let us say the employer did not give the employee the job in Round 1. In this case, both the employer and the employee earning will be 0 POINTS.

Example 2: Let us say the employer gave the employee the job in Rounds 1 and 2 but did not give the job in Round 3. In this case, the section earnings will be as follows: Round 1 earnings + Round 2 earnings.

You will be given a participation fee of 20 TL for the experiment. At the end of the experiment, the computer randomly selects 2 of the 15 completed sections. The sum of your earnings in these two selected sections determines your main stage earnings.

How will your Round earnings be determined?

The computer will randomly select one of the Green and Blue colors in each Round. The probability of choosing both colors is equal.

In each Round, the earnings or loss of the employer and employee (who is assigned to observe the computer's choice and make the decision) depend on the computer's color choice and the color decision of the employee.

Employee Decision

In each Round, the employee observes the color chosen by the computer and chooses one of the colors green and blue.

Employee:

- If Type 1; If he chooses the same color as the computer, he will benefit both himself and the employer. (employer's earnings with Type 1 employee is the same.)
- If Type 2; If he chooses the color blue, he gets a positive profit regardless of computer choice.

Round payoffs (depending on the preference of Type 1 and Type 2 employees and the computer's choice) can be observed in the tables below. The employer's earnings are the same as the Type 1 employee. If the employee chooses the same color as the computer, the employer will profit; otherwise, he will lose.

Type 1 employee's and employer's earnings

	Employee's choice = Computer's choice		Employee's choice $\neq$ Computer's choice	
	Employer's earning	Type 1 employee's earning	Employer's earning	Type 1 employee's earning
Round 1	10	10	-20	-20
Round 2	40	40	-80	-80
Round 3	160	160	-320	-320

Type 2 employee's earnings								
Computer choice	Blue				Green			
Employee choice	Blue		Green		Blue		Green	
	Employer earning	Type 2 employee earning	Employer earning	Type 2 employee earning	Employer earning	Type 2 employee earning	Employer earning	Type 2 employee earning
Round 1	10	10	-20	-20	-20	10	10	-20
Round 2	40	40	-80	-80	-80	40	40	-80
Round 3	160	160	-320	-320	-320	160	160	-320

Employer's Decision

An employer and an employee will be randomly matched when each section starts. The probability of the employee being Type 1 is 10%, probability of being Type 2 is 90%.

After the section starts, the employer decides whether to give the job to the employee he is matched with at the beginning of each Round. If the employee decides not to give the job, the section ends, and the employer is rematched with an employee for the next section.

In each section, the employer can give a maximum of 3 rounds of work to the employee with whom he is matched. The section ends when the employee completes the 3rd Round with the same employer. The employer is rematched with a new employee for the next section.

Employer:

- Decides to give or not give the job to the employee for Round 1.
 - If he gives the job, the employee makes a choice as described above. The employer sees the computer's choice and the employee's decision. They move to the 2nd Round.
 - If he does not give the job, the section ends. For the next section, the employer is randomly matched with an employee.

- Decides to give or not give the job to the employee for Round 2.
 - If he gives the job, the employee makes a choice as described above. The employer sees the computer's choice and the employee's decision. They move to the 3rd Round.
 - If he does not give the job, the section ends. For the next section, the employer is randomly matched with an employee.
- Decides to give or not give the job to the employee for Round 3.
 - If he gives the job, the employee makes a choice as described above. The employer sees the computer's choice and the employee's decision. The section ends.
 - If he does not give the job, the section ends. For the next section, the employer is randomly matched with an employee.

The employee's and employer's total section earnings are determined as the sum of their points earned from Rounds in that section. For example, suppose the employer decides to hire the employee in the first 2 Rounds and not in the 3rd Round. The points earned by both the employee and the employer from the first 2 Rounds will determine the section scores.

Stages in the Experiment.

1. Quiz. You will be asked to answer questions about the experimental design at this stage. You will get 42 POINTS for each question answered correctly.
2. The main stage. Stage involving the 15 Sections described above.
3. Extra stage. This stage is independent of the main stage. At this stage, you will be presented with 2 alternative random options. You will earn money from this section according to your decision.
4. Survey

Determination of experimental gains

You will be given a participation fee of 20 TL for the experiment.

2 of the 15 Sections will be randomly selected by the computer when the experiment is completed. Your experiment score will be the sum of points you earn from these two Sections and the points you earn from the Quiz and Extra stage. Your total experiment score will be multiplied by $1/21$ and added to the 20 TL participation fee. At the end of the experiment, you will be asked about your preferred payment type. You may receive your payment from BELIS laboratory on the day to be determined later. Alternatively, you may share your bank details so that we will send your earnings to your bank account.

D.2 Turkish Instructions and Screenshots

Giriş

Bu bir karar alma deneyidir ve bilimsel bir projenin parçasıdır.

Bu deneyin amacı, insanların farklı durumlarda nasıl karar verdiğini anlayabilmektir. Kararlarınız "doğru" ya da "yanlış" olarak değerlendirilmeyecektir.

Deneyde elde edeceğiniz kazanç, alacağınız kararlara bağlıdır ve kazancınızın ne şekilde belirleneceği bu yönergede detaylı bir şekilde açıklanmıştır. Bu nedenle, yönergeyi dikkatle okumanız ve anlamanız önemlidir. Elde edeceğiniz kazancın ödemesi, sizin tercihinize göre, BELIS laboratuvarında nakit olarak veya şahsi banka hesabınıza para transferi şeklinde yapılacaktır. Kazancınız hakkında diğer katılımcılara bilgi verilmeyecektir.

Aldığınız kararlar ve verdiğiniz cevaplar tamamen anonimdir, hiçbir kimlik bilgisi ile eşleştirilmemektedir.

Deney tamamlanıncaya kadar diğer katılımcılarla iletişim kurmanız kesinlikle yasaktır.

- ☐ Deney süresince diğer katılımcılarla iletişim kurmayacağımı onaylıyorum.

Deney süresince deney yöneticileri ile iletişiminiz sizinle uzantısı paylaşılmış olan Zoom oturumu aracılığı ile olacaktır. Deney yöneticileri bu oturum aracılığı ile sesli duyurular yapacak ve sizden gelebilecek soruları yanıtlayacaklardır. Bu nedenle, deney süresince Zoom oturumuna bağlı kalmanız ve oturuma bağlanmak için kullandığınız cihazın sesinin duyuruları işitecek şekilde açık tutmanız son derece önemlidir.

Sorularınızı Zoom oturumunun «chat» kısmını kullanarak deney yöneticilerine iletebilirsiniz.

- ☐ Deney süresince Zoom oturumunu terketmeyeceğim ve cihazın sesini duyuruları işitecek şekilde açık tutacağım

Deney yaklaşık 120 dak. sürecektir.

Deney linkleriniz kişiye özgü oluşturulmuştur.

Deney linkinizi lütfen kimseyle paylaşmayınız.

- ☐ Deney linkimi kimseyle paylaşmayacağımı onaylıyorum

Deney sırasında kopma yaşarsanız, aynı link üzerinden giriş yapıp kaldığınız yerden devam edebilirsiniz. Sorununuzun devam etmesi durumunda lütfen Zoom üzerinden deney yürütücülerinden biriyle iletişime geçin.

Deney sonlanmadan deney uygulamasını ve Zoom oturumunu terketmemeniz gerekmektedir. Tek bir katılımcı dahi deneyi tamamlamazsa, proje ekibi açısından deney verisi kullanılamaz hale gelecektir.

Deney sonlanmadan deneyi terketmeniz durumunda tarafınıza herhangi bir ödeme yapılmayacak ve bundan sonra yapılacak BELIS deneylerine katılmanız mümkün olmayacaktır.

- ☐ Deney tamamlanmadan deneyi terketmeyeceğimi onaylıyorum.

İlerle

Debug info

Basic info

ID in group	2
Group	7
Round number	1
Participant	P2
Participant label	

Giriş

Bu sayfayı tamamlamak için kalan süreniz: 9:33

Deney Yönergesi

Genel Bilgi ve Roller

Bilgisayar her Turda Yeşil ve Mavi renklerinden birini rastgele olarak seçecektir. Bilgisayarın seçtiği rengi gözlemleyemeyen bir **işveren**, bu seçimi kendisi adına gözlemleme ve buna istinaden bir karar alma işini eşleştirdiği **çalışana** verip vermeme kararı alır. İş verirse, çalışanın kararına bakarak onunla çalışmaya devam edip etmeme kararı alır.

Çalışan, 1. Tip ya da 2. Tip

- 1. Tip ise; bilgisayarla aynı rengi seçerse hem kendisine hem de işverene kazanç sağlar.
- 2. Tip ise; bilgisayarın seçiminden bağımsız olarak Mavi rengini seçerse kendisine kazanç elde eder.

İşveren, karşıdaki çalışanın 1. Tip mi yoksa 2. Tip mi olduğunu bilmez.

Deney başladığında rastgele olarak **işveren** ya da **çalışan** rolüne atanacaksınız

Eğer çalışan rolüne atanmışsanız %10 olasılıkla 1. Tip, %90 olasılıkla 2. Tip olacaksınız.

Deney sonuna kadar atandığınız rol değişmeyecektir. Yani, deney başında işveren olarak atanmışsanız deney sonlanana dek işveren, 1. Tip çalışan olarak atanmışsanız deney sonlanana dek 1. Tip çalışan, 2. Tip çalışan olarak atanmışsanız deney sonlanana dek 2. Tip çalışan olacaksınız.

Turlar ve Bölümler

Her bölüm başlarken bir işveren ile bir çalışan eşleştirilir. İşveren, eşleştirdiği çalışanın 1. Tip mi yoksa 2. Tip mi olduğunu bilmez. Sadece çalışanın 1. Tip olma olasılığının %10, 2. Tip olma olasılığının da %90 olduğunu bilir.

Her bölüm **en fazla 3 tur** içerir. Her turun başında, işveren eşleştirdiği çalışana iş verip vermeme kararı alır. İşveren, çalışana iş vermeme kararı alırsa, hangi turda olduklarından bağımsız olarak bölüm sonlanır. Bir sonraki bölüm için işveren bir çalışan ile yeniden rastgele olarak eşleştirilir.

Bölüm kazancınız nasıl belirlenecek?

Bir bölümün içinde tamamlanmış olan turlarda elde ettikleri kazanç işveren ve çalışanın bölüm kazancını belirler.

Örnek 1: Diyelim ki işveren, çalışana 1. Tur'da işi vermedi. Bu durumda hem işverenin hem de çalışanın Bölüm kazancı 0 PUAN olacaktır.

Örnek 2: Diyelim ki işveren, çalışana 1. ve 2. Tur'da işi verdi ama 3. Tur'da işi vermedi. Bu durumda Bölüm kazançları aşağıdaki şekilde olacaktır: 1. Tur kazancı + 2. Tur kazancı

Deneyde size 20TL katılım ücreti verilecektir. Deney sona erdiğinde, tamamlanmış olan 15 Bölümden 2 tanesini bilgisayar rastgele olarak seçer ve seçilen bu 2 bölümde elde etmiş olduğunuz kazançların toplamı esas aşama kazancınızı belirler.

Tur kazancınız nasıl belirlenecek?

Bilgisayar her Turda Yeşil ve Mavi renklerinden birini rastgele olarak seçecektir. Her iki rengin seçilme olasılığı eşittir.

Her Turda, işverenin ve (bilgisayarın seçimini gözlemleyip karar verme işini verdiği) çalışanın kazancı ya da kaybı **bilgisayarın renk seçimine** ve **çalışanın renk kararına** bağlıdır.

Çalışanların Kararı

Her turda, bilgisayarın seçtiği rengi gören çalışan, gördüğü rengin üzerine Yeşil veya Mavi renklerinde birini seçer.

Çalışan:

- 1. Tip ise; bilgisayarla aynı rengi seçerse hem kendisine hem de işverene kazanç sağlar. (İşverenin kazancı 1. Tip çalışanla aynıdır.)
- 2. Tip ise; bilgisayarın seçiminden bağımsız olarak Mavi rengini seçerse kendisine kazanç elde eder.

1. Tip ve 2. Tip çalışanın tercihi ve bilgisayarın tercihine göre her Turda kazançların nasıl belirleneceği aşağıdaki tablolarda görülebilir. **İşverenin kazancı 1. Tip çalışanla aynıdır.** Çalışan, bilgisayarla aynı rengi seçerse işveren kazanç sağlayacak, aksi durumda kaybedecektir.

1. Tip Çalışanın ve İşverenin Kazancı

	Çalışanın kararı = bilgisayarın seçimi		Çalışanın kararı ≠ bilgisayarın seçimi	
	İşverenin kazancı	1. Tip çalışanın kazancı	İşverenin kazancı	1. Tip çalışanın kazancı
Tur 1	10	10	-20	-20
Tur 2	40	40	-80	-80
Tur 3	160	160	-320	-320

2. Tip Çalışanın Kazancı

Bilgisayarın seçimi	Mavi				Yeşil			
	Mavi		Yeşil		Mavi		Yeşil	
Çalışanın kararı	İşverenin kazancı	2. Tip çalışanın kazancı	İşverenin kazancı	2. Tip çalışanın kazancı	İşverenin kazancı	2. Tip çalışanın kazancı	İşverenin kazancı	2. Tip çalışanın kazancı
Tur 1	10	10	-20	-20	-20	10	10	-20
Tur 2	40	40	-80	-80	-80	40	40	-80
Tur 3	160	160	-320	-320	-320	160	160	-320

İşverenin Kararı

Her Bölüm başladığında bir işveren ve bir çalışan rastgele eşleşecektir. Çalışanın 1. Tip olma ihtimali **%10**, 2. Tip olma ihtimali **%90** 'dır.

Bölüm başladıktan sonra işveren her tur başında eşleştirilmiş olduğu çalışana işi verip vermeme kararı alır. Herhangi bir turda çalışana işi vermeme kararı alırsa bölüm sonlanır ve bir sonraki bölüm için işveren yeniden bir çalışanla eşleştirilir.

Her bölümde işveren eşleştirdiği çalışana en fazla 3 tur işi verebilir. Çalışan aynı işverenle 3. turu tamamladığında bölüm sona erer ve bir sonraki bölüm için işveren yeniden bir çalışanla eşleştirilir.

İşveren:

- 1. tur için çalışana işi vermeye veya vermemeye karar verir.
 - İş verirse, çalışan yukarıda açıklandığı şekilde karar alır. İşveren bilgisayarın seçimini ve çalışanın kararını görür. 2. tura geçilir.
 - İş vermezse, bölüm sonlanır. Bir sonraki bölüm için, işveren bir çalışanla rastgele eşleştirilir.
- 2. tur için çalışana işi vermeye veya vermemeye karar verir.
 - İş verirse, çalışan yukarıda açıklandığı şekilde karar alır. İşveren bilgisayarın seçimini ve çalışanın kararını görür. 3. tura geçilir.
 - İş vermezse, bölüm sonlanır. Bir sonraki bölüm için, işveren bir çalışanla rastgele eşleştirilir.
- 3. tur için çalışana işi vermeye veya vermemeye karar verir.
 - İş verirse, çalışan yukarıda açıklandığı şekilde karar alır. İşveren bilgisayarın seçimini ve çalışanın kararını görür. Bölüm sonlanır.
 - İş vermezse, bölüm sonlanır. Bir sonraki bölüm için, işveren bir çalışanla rastgele eşleştirilir.

Çalışanın ve işverenin toplam Bölüm kazancı o bölümdeki Turlardan kazandıkları puanın toplamı olarak belirlenir. Örneğin, işveren çalışanı ilk 2 Turda işe alıp 3. Turda işe almamaya karar verirse hem çalışanın hem de işverenin ilk 2 Turdan kazandıkları puanların toplamı Bölüm puanlarını belirleyecektir.

Deney'de yer alan aşamalar.

1. **Quiz.** Bu aşamada deney tasarısı ile ilgili soruları cevaplamamız istenecektir. Doğru cevapladığınız her soru için 42 PUAN kazanacaksınız.
2. **Esas aşama.** Yukarıda açıklanan 15 Bölümden oluşan aşama.
3. **Ekstra aşama.** Bu aşama esas aşamadan bağımsızdır. Bu aşamada size 2 alternatif şansa bağlı seçenek sunulacaktır ve kararınıza göre bu bölümden para kazanacaksınız.
4. **Anket**

Deney kazançlarının belirlenmesi

Deneyde size 20TL katılım ücreti verilecektir.

Deney tamamlandığında 15 Bölümden 2'si bilgisayar tarafından rastgele olarak seçilecek ve deney puanınız bu iki Bölümden kazandığınız puan ile Quiz ve Ekstra aşamadan kazandığınız puanların toplamı olacaktır. Deney puanı **1/2 TL** ile çarpılarak 20TL katılım ücretine eklenerek tarafınıza ödenecektir.

Ödemenizi daha sonra belirlenecek gün ve saatlerde BELIS laboratuvarından alabilir veya banka hesabınıza gönderilmesi için hesap bilgilerinizi paylaşabilirsiniz. Deney sonunda hangi ödeme türünü tercih ettiğiniz sorulacaktır.

Lütfen yönergeleri inceleyin. İlerle butonu 5 dakika sonra aktif olacaktır.

Debug info

Basic info

ID in group 1

Group 7

Round number 1

Participant P1

Participant label

Session code 948f0gvvr

Sonuçlar.

Bu sayfayı tamamlamak için kalan süreniz: 4:27

5 sorudan 2 soruyu doğru cevapladınız

Bu aşamadan kazancınız: 30 PUAN

Soru	Cevabınız	Doğru cevap	Cevabınız doğru mu?
Tip 2 Çalışan, periyot kazancınızı maksimize etmek için hangisini seçmelidir?	Mavi	Mavi	Doğru
Tip 1 çalışan, bilgisayar Mavi seçtiğinde hangisini seçerse periyot kazancınızı maksimize eder?	Mavi	Mavi	Doğru
Diyeelim ki çalışanın Tip 1 olma ihtimalini 60% olarak seçtiniz. Aşağıdakilerden hangisi doğrudur?	Eğer çalışan Tip 1 ise, $(1-60/100)^2 = 0.16$ olasılıkla 2 PUAN , $1-(1-60/100)^2 = 0.84$ olasılıkla ise 0 PUAN kazanacaksınız.	Eğer çalışan Tip 1 ise, $1 - (1-60/100)^2 = 0.84$ olasılıkla 2 PUAN , $(1-60/100)^2 = 0.16$ olasılıkla ise 0 PUAN kazanacaksınız.	Yanlış
İfadenin doğru mu yanlış mı olduğunu seçin: Deney süresince rolünüz (işveren veya çalışan) ve tipiniz (Tip 1 veya Tip 2 çalışan) sabit kalacaktır.	Doğru	Deney başında bazı katılımcılar rastgele olarak "Çalışan" rolünde, bazı katılımcılar "İşveren" rolünde olacak. Rolünüz deney boyunca sabit kalacaktır. Her Bölüm başında çalışanlara bilgisayar bir "Tip" atayacak ve çalışanlar 0,1 ihtimalle Tip 1, 0,9 ihtimalle Tip 2 olacaklardır. Tipiniz Bölüm boyunca sabit kalacak, her bölüm bittiğinde tekrar rastgele bir tip atanacaktır.	Yanlış
İfadenin doğru mu yanlış mı olduğunu seçin: İşverenin 1. Tip ve 2. Tip çalışan ile eşleşme ihtimali eşittir.	Doğru	İşveren 0,1 ihtimalle Tip 1, 0,9 ihtimalle Tip 2 çalışanla eşleşecektir.	Yanlış

Lütfen sayfayı inceleyin. İlerle butonu 2 dakika sonra aktif olacaktır.

Debug info

Basic info

ID in group	1
Group	3
Round number	1
Participant	P1
Participant label	
Session code	t19aao5a

Bu sayfayı tamamlamak için kalan süreniz: **0:06**

İşveren rolüne atandınız. Rolünüz deney boyunca sabit kalacaktır.

İlerle

Bölüm : 1 1. Tur İş verme kararı.

İŞVEREN rolündesiniz. **Bir çalışanla rastgele olarak eşleştirildiniz.** Çalışan 10% olasılık ile 1. Tip, 90% olasılık ile 2. Tip.

Çalışanın yüzde kaç ihtimalle 1. Tip olduğunu düşünüyorsunuz?

Çalışana 1. Tur için işi vermek istiyor musunuz?

- ☒ İş ver
☐ İş vermem

Çalışana işi verirsiniz Tur kazancınız; Çalışanın kararı bilgisayarın seçimiyle aynı olursa 10 PUAN, bilgisayarın seçiminden farklı olursa -20 PUAN olacak.

İlerle

Bölüm : 1

1. Tur Çalışan Kararı.

1. Tip çalışan rolündesiniz.

İşveren 1. Tur için size işi verdi.

Detaylı kazanç şeması için aşağıdaki yönergelere bakınız.

Bilgisayarın Tercihi: **Mavi**

Kararınız :

☐ Yeşil ☒ Mavi

Bilgisayarın seçimi Mavi iken Mavi seçerseniz sizin Tur kazancınız 10 PUAN, İşverenin Tur kazancı 10 PUAN olacak.

İlerle

Bölüm : 1

1. Tur Sonuçları

Bu sayfayı tamamlamak için kalan süreniz: **0:17**

Geçmiş Tur Kararları

	Bilgisayarın seçimi	Çalışanın kararı	Kazancınız	Tip Tahmininiz
1. Tur	Mavi	Mavi	10 PUAN	23

Bilgisayarın seçimi **Mavi** iken çalışanın kararı **Mavi** oldu.Bu yüzden 1. Tur kazancınız: **10 PUAN**.

İlerle

Bölüm : 1

2. Tur İş verme kararı.

İŞVEREN rolündesiniz.

Geçmiş Tur Kararları				
	Bilgisayarın seçimi	Çalışanın kararı	Kazancınız	Tip Tahmininiz
1. Tur	Mavi	Mavi	10 PUAN	23

Çalışanın yüzde kaç ihtimalle 1. Tip olduğunu düşünüyorsunuz?

Çalışana 2. Tur için işi vermek istiyor musunuz?

- ☒ İş ver
☐ İş verme

Çalışana işi verirsiniz Tur kazancınız: Çalışanın kararı bilgisayarın seçimiyle aynı olursa 40 PUAN, bilgisayarın seçiminden farklı olursa -80 PUAN olacak.

İlerle

Bölüm : 1

2. Tur Çalışan Kararı.

1. Tip çalışan rolündesiniz.

İşveren size 2. Tur için işi verdi.

Detaylı kazanç şeması için aşağıdaki yönergelere bakınız.

Geçmiş Tur Kararları			
	Bilgisayarın seçimi	Kararınız	Kazancınız
1. Tur	Mavi	Mavi	10 PUAN

Bilgisayarın Seçimi: **Yeşil**

Kararınız :

- ☒ Yeşil ☐ Mavi

Bilgisayarın tercihi Yeşil iken Yeşil seçerseniz sizim Tur kazancınız 40 PUAN, İşverenin Tur kazancı 40 PUAN olacak.

İlerle

Bölüm : 1

2. Tur Sonuçları

Bu sayfayı tamamlamak için kalan süreniz: **0:16**

Geçmiş Tur Kararları				
	Bilgisayarın seçimi	Çalışanın kararı	Kazancınız	Tip Tahmininiz
1. Tur	Mavi	Mavi	10 PUAN	23
2. Tur	Yeşil	Yeşil	40 PUAN	12

Bilgisayarın seçimi **Yeşil** iken çalışanın kararı **Yeşil** oldu.

Bu yüzden 2. Tur kazancınız: **40 PUAN**.

İlerle

Bölüm : 1

3. Tur İş verme kararı.

İŞVEREN rolündesiniz.

Geçmiş Tur Kararları				
	Bilgisayarın seçimi	Çalışanın kararı	Kazancınız	Tip Tahmininiz
1. Tur	Mavi	Mavi	10 PUAN	23
2. Tur	Yeşil	Yeşil	40 PUAN	12

Çalışanın yüzde kaç ihtimalle 1. Tip olduğunu düşünüyorsunuz?

12

Çalışana 3. Tur için işi vermek istiyor musunuz?

- ☐ İş ver
☒ İş verme

Çalışana işi vermezseniz, Bölüm sonlanacak. Bu Turdan 0 PUAN kazanacaksınız.

İlerle

Bölüm : 1

3. Tur Çalışan Kararı.

1. Tip çalışan rolündesiniz.
İşveren size 3. Tur için işi verdi.

Detaylı kazanç şeması için aşağıdaki yönergelere bakınız.

Geçmiş Tur Kararları			
	Bilgisayarın seçimi	Kararınız	Kazancınız
1. Tur	Mavi	Mavi	10 PUAN
2. Tur	Yeşil	Yeşil	40 PUAN

Bilgisayarın Tercihi: **Mavi**

Kararınız :

☐ Yeşil ☒ Mavi

Bilgisayarın seçimi Mavi iken Mavi seçerseniz sizin Tur kazancınız 160 PUAN, İşverenin Tur kazancı 160 PUAN olacak.

İlerle

Bölüm : 1

3. Tur Sonuçları

Bu sayfayı tamamlamak için kalan süreniz: **0:17**

Geçmiş Tur Kararları			
	Bilgisayarın seçimi	Kararınız	Kazancınız
1.Tur	Mavi	Mavi	10 PUAN
2.Tur	Yeşil	Yeşil	40 PUAN
3.Tur	Mavi	Mavi	160 PUAN

Bilgisayarın seçimi **Mavi** iken kararınız **Mavi** oldu.

Bu yüzden 3.Tur kazancınız: **160 PUAN**.

İlerle

Bölüm 1 Sonuçları

Bu sayfayı tamamlamak için kalan süreniz: **0:57**

Geçmiş Tur Kararları

	Bilgisayarın seçimi	Çalışanın kararı	Kazancınız	Tip Tahmininiz
1.Tur	Mavi	Mavi	10 PUAN	23
2.Tur	Yeşil	Yeşil	40 PUAN	12
3.Tur	Mavi	Mavi	160 PUAN	12

1. Tur kazancınız: **10 PUAN**.

2. Tur kazancınız: **40 PUAN**.

3. Tur kazancınız: **160 PUAN**.

Toplam Bölüm kazancınız **210 PUAN**.

İlerle

Lütfen diğer katılımcıları bekleyin

Rastgele bir katılımcı ile yeniden eşleştirileceksiniz. Lütfen Bekleyin.

Bölüm 15 Sonuçları

Bu sayfayı tamamlamak için kalan süreniz: **0:56**

Geçmiş Tur Kararları

	Bilgisayarın seçimi	Kararınız	Kazancınız
1.Tur	Mavi	İş verilmedi	0 PUAN
2.Tur	Mavi	İş verilmedi	0 PUAN
3.Tur	Mavi	İş verilmedi	0 PUAN

İşveren size 1. Turda işi vermedi.

Toplam Bölüm kazancınız **0 PUAN**.

İlerle

Son

Bu sayfayı tamamlamak için kalan süreniz: **0:49**

Ödeme için seçilen bölümler: 6 ve 13
Bu Bölümlerden kazancınız: 0 ve 0 PUAN.

Birazdan Ekstra Aşamaya yönlendirileceksiniz.

İlerle

Ekstra Aşama

Bu sayfayı tamamlamak için kalan süreniz: **3:58**

Şimdi, 10 adet tercih yapmanız istenecektir. Her tercih "Seçenek A" ve "Seçenek B" arasında yapılacaktır. Bu seçeneklerdeki olası kazançlar sabit olsa da yüksek kazancın olasılığı değişecektir.

Bütün kararları verdikten sonra, 10 karardan biri rastgele olarak seçilecek ve bu aşamadaki kazancınızı belirleyecektir. Bu kararda, seçmiş olduğunu seçeneğe göre (A veya B) yüksek kazanç mı düşük kazanç mı kazandığınızı rastgele olarak (belirtilen olasılıklara göre) belirlenecektir.

Özet olarak: 10 adet tercih yapacaksınız; her tercih için "Seçenek A" veya "Seçenek B"yi seçeceksiniz. Bazı satırlarda 'A', diğerlerinde 'B' seçebilirsiniz. Tercihlerinizi bitirdiğinizde 10 karardan biri rastgele olarak seçilecek. O karardaki kazancınız belirtilen olasılıklara göre rastgele olarak belirlenecektir.

İlerle

Bu sayfayı tamamlamak için kalan süreniz: **1:23**

	Seçenek A		Seçenek B	
	10.00% olasılık ile 250 PUAN, yoksa 190 PUAN	☒ ☐	10.00% olasılık ile 420 PUAN, yoksa 60 PUAN	
	20.00% olasılık ile 250 PUAN, yoksa 190 PUAN	☒ ☐	20.00% olasılık ile 420 PUAN, yoksa 60 PUAN	
	30.00% olasılık ile 250 PUAN, yoksa 190 PUAN	☐ ☒	30.00% olasılık ile 420 PUAN, yoksa 60 PUAN	
	40.00% olasılık ile 250 PUAN, yoksa 190 PUAN	☐ ☒	40.00% olasılık ile 420 PUAN, yoksa 60 PUAN	
	50.00% olasılık ile 250 PUAN, yoksa 190 PUAN	☐ ☒	50.00% olasılık ile 420 PUAN, yoksa 60 PUAN	
	60.00% olasılık ile 250 PUAN, yoksa 190 PUAN	☐ ☒	60.00% olasılık ile 420 PUAN, yoksa 60 PUAN	
	70.00% olasılık ile 250 PUAN, yoksa 190 PUAN	☐ ☒	70.00% olasılık ile 420 PUAN, yoksa 60 PUAN	
	80.00% olasılık ile 250 PUAN, yoksa 190 PUAN	☐ ☒	80.00% olasılık ile 420 PUAN, yoksa 60 PUAN	
	90.00% olasılık ile 250 PUAN, yoksa 190 PUAN	☐ ☒	90.00% olasılık ile 420 PUAN, yoksa 60 PUAN	
	100.00% olasılık ile 250 PUAN, yoksa 190 PUAN	☐ ☒	100.00% olasılık ile 420 PUAN, yoksa 60 PUAN	

İlerle

Ekstra Aşama Sonuçları

Aşağıdaki karar rastgele seçilmiştir:

	Option A		Option B	
	30.00% olasılıkla 250 PUAN , yoksa 190 PUAN	☐	30.00% olasılıkla 420 PUAN , yoksa 60 PUAN	

Yukarıda görüleceği gibi bu kararda B seçtiniz. Tercihinize göre, belirtilen kazancınız belirtilen olasılıklar doğrultusunda rastgele belirlenmiştir.

Bu aşamadan kazancınız: **60 PUAN**.

İlerle

Anket

Anket sorularını dikkatli bir şekilde cevaplamanız önemlidir. Ankete katılımınız için deneyde elde etmiş olduğunuz tüm kazançta 10 TL daha eklenecektir.

İlerle

Anket

Lütfen aşağıdaki soruları cevaplayın.

Genel olarak ne kadar risk alırsınız? 0 = kesinlikle risk almam, 10 = kesinlikle risk alırım.

☐ 0 ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10

Aşağıdaki alternatiflerden hangisini tercih edersiniz?

- ☐ Kesin olarak 475 TL kazanacaksınız
- ☐ %25 ihtimalle 2000TL kazanacaksınız, %75 ihtimalle kazancınız olmayacak

Aşağıdaki alternatiflerden hangisini tercih edersiniz?

- ☐ Kesin olarak 725 TL kaybedeceksiniz
- ☐ %75 ihtimalle 1000TL kaybedeceksiniz, %25 ihtimalle kaybınız olmayacak

İlerle

Anket

Lütfen aşağıdaki soruları cevaplayın.

Yaşınız:

Cinsiyetiniz:

- ☐ Erkek
- ☐ Kadın

BİLGİ ye giriş yılınız:

Bölümünüz:

ÖSYM başarı Burs durumunuz:

- ☐ Tam Burslu
- ☐ %75 Burslu
- ☐ %50 Burslu
- ☐ %25 Burslu
- ☐ Burssuz

Babanızın son mezun olduđu eğitim kurumunu hangisi en iyi ifade eder?

- ☐ Hiç okula gitmemiş
- ☐ İlk veya orta öğretim
- ☐ Lise
- ☐ Ön lisans / Meslek Yüksek okulu
- ☐ Üniversite (Lisans)
- ☐ Yüksek Lisans
- ☐ Doktora

Annenizin son mezun olduđu eğitim kurumunu hangisi en iyi ifade eder?

- ☐ Hiç okula gitmemiş
- ☐ İlk veya orta öğretim
- ☐ Lise
- ☐ Ön lisans / Meslek Yüksek okulu
- ☐ Üniversite (Lisans)
- ☐ Yüksek Lisans
- ☐ Doktora

Daha önce bir BELİS deneyine katıldınız mı?

- ☐ Evet
- ☐ Hayır

Deneydeki kararlarınızı nasıl aldığınızı açıklayınız.

Deneydeki kararlarınızı en iyi açıklayan seçeneği seçiniz. :

- ☐ İlk iki turda bilgisayar Yeşil seçmişse 3. Turda Yeşil seçme olasılığının %50 olduğunu düşündüm.
- ☐ İlk iki turda bilgisayar Yeşil seçmişse 3. Turda Yeşil seçme olasılığının %40-%50 arasında olduğunu düşündüm.
- ☐ İlk iki turda bilgisayar Yeşil seçmişse 3. Turda Yeşil seçme olasılığının %30-%40 arasında olduğunu düşündüm.
- ☐ İlk iki turda bilgisayar Yeşil seçmişse 3. Turda Yeşil seçme olasılığının %20-%30 arasında olduğunu düşündüm.
- ☐ İlk iki turda bilgisayar Yeşil seçmişse 3. Turda Yeşil seçme olasılığının %20den daha az olduğunu düşündüm.

Deneydeki kararlarınızı en iyi açıklayan seçeneği seçiniz. :

- ☐ Herhangi bir turda bilgisayar Yeşil seçmişse bir sonraki turda Yeşil seçme olasılığının %50 olduğunu düşündüm.
- ☐ Herhangi bir turda bilgisayar Yeşil seçmişse bir sonraki turda Yeşil seçme olasılığının %40-%50 arasında olduğunu düşündüm.
- ☐ Herhangi bir turda bilgisayar Yeşil seçmişse bir sonraki turda Yeşil seçme olasılığının %30-%40 arasında olduğunu düşündüm.
- ☐ Herhangi bir turda bilgisayar Yeşil seçmişse bir sonraki turda Yeşil seçme olasılığının %20-%30 arasında olduğunu düşündüm.
- ☐ Herhangi bir turda bilgisayar Yeşil seçmişse bir sonraki turda Yeşil seçme olasılığının %20den daha az olduğunu düşündüm.

İlerle

Deneyimiz sona erdi.

Kazancınız:

52,00 TL

Lütfen aşağıdaki bilgileri kaydedin:

Tarih: **2022-06-04**

Katılımcı numarası: **1**

Katılımcı kodu: **f1gmz81c**

Tercih ettiğiniz ödeme yöntemini seçiniz.

- ☐ Deney kazancımın aşağıda bildireceğim hesaba gönderilmesini onaylıyorum.
- ☒ Deney kazancımı BELİS laboratuvarından almak istiyorum

Seçiminiz kaydedildi. Sizinle birkaç gün içinde ödemenin yapılacağı zamanla ilgili e posta yoluyla iletişime geçilecektir.

İlerle

Tercih ettiğiniz ödeme yöntemini seçiniz.

- ☒ Deney kazancımın aşağıda bildireceğim hesaba gönderilmesini onaylıyorum.
- ☐ Deney kazancımı BELİS laboratuvarından almak istiyorum

Deney ödemeniz projenin 2021 mali yılı bütçesinden yapılacağı için 2021'in ilk iş gününde (4 Ocak 2021) yapılacaktır. Ödemenizi daha erken bir tarihte talep ediyorsanız lütfen mail yoluyla belirtin.

Deney ödemenizin banka hesabına gönderilmesini talep ettiyseniz devam etmeden önce aşağıdaki adımları izleyin.

- 1 - Aşağıdaki formda yer alan ifadeleri kendi el yazınızla boş bir A4 kâğıdı üzerine yazın. Lütfen formdaki tüm ifadeleri eksiksiz ve okunaklı şekilde yazdığınızdan emin olunuz.
- 2- Doldurmuş olduğunuz formun net bir fotoğrafını çekip belis@bilgi.edu.tr adresine yukarıdaki tarih, katılımcı numarası ve katılımcı kodu bilgileri ile birlikte gönderin.
- 3- Üniversitemizin Mali İşler biriminden gelebilecek herhangi bir ilave talebe istinaden formun aslını birkaç ay saklamanız rica olunur.
- 4- Sayfanın aşağısındaki bilgileri eksiksiz olarak doldurun.

İstanbul Bilgi Üniversitesi Ekonomi Laboratuvarı (BELİS)'te yapılan deneylere katılımım karşılığında kazandığım TL (yazı ileTL) tutarın, belirttiğim banka hesabına transfer edilmesini kabul ettiğimi beyan ederim.

Ad-Soyad:

TC Kimlik no:

IBAN:

İmza:

Appendix E

CHAPTER 3 PROOFS

I first define parameter $\delta_i \in (0, 1)$ as the relative weight of the future periods while $1 - \delta_i$ is the weight of period i . In other words, in each period i the proportion of remaining stakes allocated to period i is $(1 - \delta_i)$ and the proportion to leave to subsequent periods is δ_i . Thus, $\gamma_i = (1 - \delta_i) \prod_{j=1}^{i-1} \delta_j$. Defining the relative stakes this way allows us to set stakes dynamically.

Proof. **Proposition 3.1**

1. Suppose $\rho_{i+1}(0, 0) \leq \rho_{i+1}(0, 1)$ for some i in the principal-optimal equilibrium.

Case 1. The agent prefers lower reputation level in period $i+1$ (i.e., lower reputation level leads to a higher payoff). Then, given $\theta_i = 0$, the good agent prefers $a_i = 0$ as this action leads to a higher period i payoff and period $i+1$ continuation payoff. But then, $\mu_i^0 = 1$ implies $\rho_{i+1}(0, 0) \geq \rho_{i+1}(0, 1)$, contradiction.

Case 2. The agent prefers higher reputation level:

1. $\nu_i^0 = 0$ as the bad agent has no incentive to play action 0. If $\mu_i^0 > 0 \implies \rho_{i+1}(0, 0) = 1$, contradiction. Hence, it is only possible if $\mu_i^0 = 0$, i.e., the agents pool given $\theta_i = 0$ and $\rho_{i+1}(0, 1) = \rho_i$. If the of world is 0, then the principal gets a payoff of $k - \frac{1}{2}$ in period i .

Now, suppose the above defined pooling strategies constitute an equilibrium and consider the following: $\mu_i^0 = 1 = \nu_i^0$. Unlike above, $\rho_{i+1}(0, 0) = \rho_i$. Markovian property implies that the equilibrium play and hence the payoffs in period $i+1$ onward is the same irrespective of the action on which the agents pool. When the agents pool on action 0, the principal's period i payoff is k . The principal prefers the agents to pool on action 0. Therefore, $\mu_i^0 = 0$ can not hold in the principal-optimal equilibrium.

2. $\mu_i^0 = 1 \implies \rho_{i+1}(0, 0) > \rho_{i+1}(0, 1)$, contradiction.

3. $\mu_i^0 \in (0, 1)$, i.e., the bad agent plays both actions with a positive probability. As higher reputation level is preferred, the good agent will play action 1 with probability 1. But then, $\rho_{i+1}(0, 0) = 1$, contradiction.

Hence, $\rho_{i+1}(0, 0) > \rho_{i+1}(1, 0)$ for all i in the principal-optimal equilibrium.

2. Next, we define reputation incentives given $\theta_i = 1$. We show that playing action 1 given $\theta_i = 1$ can never increase the agent's reputation level. Suppose for contradiction that $\rho_i < \rho_{i+1}(1, 1)$ for some i in the principal-optimal equilibrium. First, note that the bad agent need to play action 0 with a positive probability for $\rho_i < \rho_{i+1}(1, 1)$ to hold. Suppose not:

1. If $\nu_i^1 = 0$ and $\mu_i^1 = 0$, then $\rho_i = \rho_{i+1}(1, 1)$, contradiction.
2. If $\nu_i^1 = 0$ and $\mu_i^1 > 0$, then $\rho_{i+1}(0, 1) = 1 > \rho_{i+1}(1, 1) = 0$, another contradiction.

Hence, we need to have $\nu_i^1 > 0$. Consider the following cases.

Case 1. The agent prefers higher reputation level in period $i + 1$ (i.e., lower reputation level leads to a higher payoff). Then, given $\theta_i = 1$, the bad agent plays $a_i = 1$ with a probability 1 as this action leads to a higher period i payoff and period $i + 1$ continuation payoff. If $\mu_i^1 > 0$ then $\rho_{i+1}(1, 0) = 1 > \rho_i > \rho_{i+1}(1, 1)$, contradiction. If $\mu_i^1 = 0$ then $\rho_i = \rho_{i+1}(1, 1)$, another contradiction.

Case 2. The agent prefers higher reputation level in period $i + 1$. Note that, the future expected payoff of the good agent is strictly higher than that of the bad agent, i.e., $V_{i+1}^g > V_{i+1}^b$. The reason is as follows. The good agent has a $1/2$ chance of taking the action that will increase her reputation (playing action 0 when $\theta = 0$). Playing action 0 in any period costs the bad agent that periods payoff, but the good agent gets the full payoff of k .

If $\nu_i^1 > 0$, then $\mu_i^1 = 1$ as the good agent values future more. But then, $\rho_{i+1}(1, 0) > \rho_{i+1}(1, 1) = 0$, contradiction.

□

Remark 1. *There are cases where the principal is indifferent among two possible equilibrium behavior: (1) pooling when $\theta_i = 1$ and (2) reputation incentive when $\theta_i = 1$. In the first case, both types of agent play action 1 when $\theta_i = 1$, i.e., there is no reputation update. The principal's out-of-equilibrium path belief $\rho_{i+1}(0, 1)$ is sufficiently low so that the agent has no incentive to play action 0. This behavior is observed in the experiment in almost all instances; hence we will define the babbling strategy in such cases. In the experiments, good agents always preferred to play the action matching the state of world. Hence we focus on the equilibrium where the agent's pool given $\theta_i = 1$. More general equilibrium strategies are detailed in the proof of each proposition (if applicable).*

Consider the reputation building behavior of the agent. If the agents pool in period i when $\theta_i = 1$, then the observation $\{a_i = 0, \theta_i = 1\}$ is out-of-equilibrium path. Hence, the principal's belief $\rho_{i+1}(0, 1)$ is free - as long as it is sufficiently low so that the agent's sequential rationality constraint is satisfied. In such case, $\rho_{i+1}(1, 1) = \rho_i > \rho_{i+1}(0, 1)$. On the other hand, if the agents separate given $\theta_i = 1$, then $\rho_{i+1}(0, 1) > \rho_i > \rho_{i+1}(1, 1)$ must hold. In words, playing action 1 when the state of world is 1 does not increase the agent's reputation level in any case.

Proof. Proposition 3.2

I will provide the proof by backward induction. First, we will provide two periods equilibrium.

Lemma E.1. *In two-periods interaction, the principal-optimal equilibrium:*

1. *If $\rho_1 \geq 1 - 2k$:*

$$\delta_1 = \frac{1}{1+k}$$

$$\nu_1^1 = \mu_1^1 = 0 \ \& \ \nu_1^0 = 1$$

$$\lambda_2(0, 0) = \lambda_2(1, 1) = 1 \ \& \ \lambda_2(1, 0) = \lambda_2(0, 1) = 0$$

2. *If $1 - 4k + 4k^2 \leq \rho_1 < 1 - 2k$:*

$$\delta_1 = \frac{1}{1+k}$$

$$\nu_1^1 = \mu_1^1 = 0 \text{ \& } \nu_1^0 = \frac{2k\rho_1}{(1-\rho_1)(1-2k)}$$

$$\lambda_2(0,0) = 1, \text{ \& } \lambda_2(1,1) = \lambda_2(1,0) = \lambda_2(0,1) = 0$$

Proof. Lemma E.1

The expected continuation payoff of the principal in the last period ($i = 2$) is given by:

$$V_2^P = k - \frac{1}{2} + \frac{1}{2}\rho_2$$

Hence, $\lambda_2(\theta_1, a_1) = 1$ iff $\rho_2(\theta_1, a_1) \geq 1 - 2k$

1. $\rho_1 \geq 1 - 2k$:

First note that $\nu_1^1 = \mu_1^1 = 0$ in the principal-optimal equilibrium. In any equilibrium where $\nu_1^1 > 0$ or $\mu_1^1 > 0$, the principal gets a lesser payoff.

We will show that the principal chooses $\delta_1 = \frac{1}{1+k}$. For contradiction, suppose that $\delta_1 > \frac{1}{1+k}$. Then, $\nu_1^0 = 1$ and $V_1^P = (1 - \delta_1)k + \delta_1(k - \frac{1}{2} + \frac{1}{2}\rho_1)$ which is decreasing in δ_1 . Hence, the principal will deviate to $\delta_1 = \frac{1}{1+k}$. Suppose $\delta_1 < \frac{1}{1+k}$, then $\nu_1^0 = 0$ and

$$V_1^P = (1 - \delta_1) \left(\rho_1 k + (1 - \rho_1) \left(k - \frac{1}{2} \right) \right) + \underbrace{\frac{1}{2} \delta_1 \rho_1 k}_{\theta_1=0} + \underbrace{\frac{1}{2} \delta_1 \left(k - \frac{1}{2} + \frac{1}{2}\rho_1 \right)}_{\theta_1=1}$$

Which is increasing in δ_1 . Hence, the principal will deviate to $\delta_1 = \frac{1}{1+k}$. Then:

$$\begin{aligned} V_1^P &= \frac{k}{1+k} \left(\underbrace{\rho_1 k + \frac{1}{2}(1-\rho_1)(k-1)}_{\theta_1=1 \& b} + \underbrace{\frac{1}{2}(1-\rho_1)(k-1+\nu_1^0)}_{\theta_1=0 \& b} \right) \\ &+ \underbrace{\frac{1}{2} \frac{1}{1+k}}_{\theta_2=0} \underbrace{(\rho_1 + (1-\rho_1)\nu_1^0)}_{\text{prob}[a_1=0|\theta_1=0]} \left(k - \frac{1}{2} + \frac{1}{2}\rho_2(0,0) \right) \\ &+ \underbrace{\frac{1}{2} \frac{1}{1+k}}_{\theta_2=1} \left(k - \frac{1}{2} + \frac{1}{2}\rho_1 \right) \end{aligned}$$

For $k > \frac{1}{4}$, $\frac{\partial V_1^P}{\partial \nu_1^0} = \frac{1}{2} \frac{1}{1+k} (1 - \rho_1) (2k - \frac{1}{2}) > 0$. Hence, V_1^P is increasing in ν_1^0 and the principal-optimal $\nu_1^0 = 1$. $\rho_2(0, 0) = \rho_2(1, 1) = \rho_1 \geq 1 - 2k$ & $\rho_2(1, 0) = \rho_2(0, 1) = 0$, thus $\lambda_2(0, 0) = \lambda_2(1, 1) = 1$ & $\lambda_2(1, 0) = \lambda_2(0, 1) = 0$.

V_1^P becomes:

$$V_1^P = \frac{1}{1+k} \left(k^2 + k - \frac{1}{2} + \frac{1}{2} \rho_1 \right)$$

2. $(1 - 2k)^2 \leq \rho_1 < 1 - 2k$:

Note that $\lambda_2(\theta_1, a_1) = 1$ can not be an equilibrium strategy because it implies $\nu_1^{\theta_1} = 0$ and $\rho_2(\theta_1, 1) \leq \rho_1 < 1 - 2k$ contradiction.

Principal again chooses $\delta_1 = \frac{1}{1+k}$. For contradiction, suppose that $\delta_1 > \frac{1}{1+k}$. Then,

1. $\lambda_2(\theta_1, 0) = 1$ for some $\theta_1 \in \{0, 1\}$. Sequential rationality of bad agent implies $\nu_1^{\theta_1} = 1$. But then, $\rho_2(\theta_1, 0) = \rho_1 < 1 - 2k$, contradicting $\lambda_2(\theta_1, 0) = 1$.
2. Hence, $\lambda_2(\theta_1, 0) < 1$ for both $\theta_1 \in \{0, 1\}$. $\lambda_2(\theta_1, 0) < 1$ implies $V_2^P = 0$. As $V_2^P = 0$, V_1^P is decreasing in δ_1 , hence the principal will deviate to $\delta_1 = \frac{1}{1+k}$.

Now suppose suppose that $\delta_1 < \frac{1}{1+k}$:

$$V_1^P = (1 - \delta_1) \left(\rho_1 k + (1 - \rho_1) \left(k - \frac{1}{2} \right) \right) + \underbrace{\frac{1}{2} \delta_1 \rho_1 k}_{\theta_1=0}$$

Which is increasing in δ_1 . Hence, the principal will deviate to $\delta_1 = \frac{1}{1+k}$.

For $\delta_1 = \frac{1}{1+k}$, the expected continuation payoff of the principal:

$$\begin{aligned} V_1^P &= \frac{k}{1+k} \left(\frac{1}{2} \rho_1 k + \frac{1}{2} \rho_1 (k - \mu_1) + \frac{1}{2} (1 - \rho_1) (k - \nu_1^1) + \frac{1}{2} (1 - \rho_1) (k - 1 + \nu_1^0) \right) \\ &\quad + \underbrace{\frac{1}{2} \frac{1}{1+k}}_{\theta_2=0} \underbrace{(\rho_1 + (1 - \rho_1) \nu_1^0)}_{\text{prob}[a_1=0|\theta_1=0]} \left(k - \frac{1}{2} + \frac{1}{2} \rho_2(0, 0) \right) \\ &\quad + \underbrace{\frac{1}{2} \frac{1}{1+k}}_{\theta_2=1} \underbrace{(\rho_1 \mu_1 + (1 - \rho_1) \nu_1^1)}_{\text{prob}[a_1=0|\theta_1=1]} \left(k - \frac{1}{2} + \frac{1}{2} \rho_2(0, 1) \right) \end{aligned}$$

$\frac{\partial V_1^P}{\partial \nu_1^1} = -\frac{1}{2} \frac{1}{1+k} (1 - \rho_1) < 0$, hence V_1^P is decreasing in ν_1^1 and the principal-optimal $\nu_1^1 = 0$.

V_1^P is independent of μ_1 , thus the principal-optimal $\mu_1 \in [0, 1]$. We choose $\mu_1 = 0$, i.e., the babbling behavior given $\theta_1 = 1$ as our equilibrium behavior.

$\frac{\partial V_1^P}{\partial \nu_1^0} = \frac{1}{2} \frac{1}{1+k} (1 - \rho_1) (2k - \frac{1}{2}) > 0$ for $k > \frac{1}{4}$, hence V_1^P is increasing in ν_1^0 . The principal-optimal ν_1^0 is the maximum probability in which the equilibrium believes are satisfied. More precisely, $\rho_2(0, 0) \geq 1 - 2k$ so that $\lambda_2(0, 0) = 1$. $\rho_2(0, 0)$ is decreasing in ν_1^0 and the maximum ν_1^0 satisfying equilibrium conditions is

$$\frac{2k\rho_1}{(1 - \rho_1)(1 - 2k)} \text{ and } \rho_2(0, 0) = 1 - 2k.$$

V_1^P becomes:

$$V_1^P = \frac{k}{(1 + k)(1 - 2k)} \left(2k - 2k^2 - \frac{1}{2} + \frac{1}{2}\rho_1 \right)$$

Which is non-negative for $\rho_1 \geq (1 - 2k)^2$. Thus, for lower initial reputation levels, $\lambda_1 = 0$

□

Proof. **Proposition 3.2**

From Lemma E.1, we know last two periods equilibria for all possible ranges of ρ_2 . Note that, by Proposition 3.1, $\rho_{i+1}(\theta_i, 1) \leq \rho_i$ for $\theta_i \in \{0, 1\}$. As $\rho_1 < (1 - 2k)^2$, $\lambda_2(\theta_1, 0) = 0$ in any equilibrium. In the first period of the three periods interaction, there are two possibilities: (1) $\rho_2(\theta_1, 0) \geq 1 - 2k$ or (2) $\rho_2(\theta_1, 0) < 1 - 2k$. If $a_1 = 0$ is played with a positive probability, then both $\rho_2(\theta_1, 0)$ and $\rho_2(\theta_1, 0)$ should belong to the same range. Otherwise, suppose $\rho_2(\theta'_1, 0) \geq 1 - 2k$ and $\rho_2(\theta''_1, 0) < 1 - 2k$ for $\theta'_1 \neq \theta''_1$. For δ_1 , if the bad agent is indifferent between both actions given $\theta_1 = \theta'_1$, then $\nu_1^{\theta''_1} = 1$ as $(\theta''_1, 0)$ leads to a higher continuation payoff than $(\theta'_1, 0)$. It implies $\rho_2(\theta''_1, 0) = \rho_1 < 1 - 2k$, contradiction. For δ_1 , if the bad agent is indifferent between both actions given $\theta_1 = \theta''_1$, then $\nu_1^{\theta'_1} = 0$. It implies $\rho_2(\theta'_1, 0) = 1 > 1 - 2k$, contradiction.

Following the similar calculations as in Lemma E.1, it is easy to show that V_1^P :

- is increasing in δ_1 if δ_1 is such that $\nu_1^{\theta_1} = 0$
- is decreasing in δ_1 if δ_1 is such that $\nu_1^{\theta_1} \geq 0$

hence, the principal will set δ_1 so that the bad agent is indifferent between playing both actions.

Next, we define the 4 different equilibrium and select the principal-optimal one. Following Proposition 3.1 and Lemma E.1, we know that the agent is not hired in the second period if action 1 is observed in the first period. Hence, the only possible action after which the agent is hired in the second period is action 0. There are two cases: $\rho_2(0,0) \geq 1 - 2k$ or $\rho_2(0,0) < 1 - 2k$. In each case, there are two sub-cases: agents pool when $\theta_1 = 1$ or they separate, i.e., at least one type plays action 0 with a positive probability.

Case 1. $\rho_2(0,0) \geq 1 - 2k$.

Case 1.1. $\mu_1 > 0$.

The sequential rationality constraint of the bad agent implies:

$$\begin{aligned} (1 - \delta_1)k &= (1 - \delta_1)(k - 1) + \delta_1 \left(\frac{1}{2}k + \frac{k^2}{1 + k} \right) \\ \Leftrightarrow \delta_1 &= \frac{2(1 + k)}{2 + 3k + 3k^2} \end{aligned}$$

For given δ_1 , $\mu_1 = 1$ as $V_2^g > V_2^b$. It is easy to show that V_1^P is decreasing in ν_1^1 . Hence the principal optimal $\nu_1^1 = 0$. The expected continuation payoff of the principal:

$$\begin{aligned} V_1^P &= \frac{k + 3k^2}{2 + 3k + 3k^2} \left(\rho_1 \left(k - \frac{1}{2} \right) + (1 - \rho_1) \left(\frac{1}{2}k + \frac{1}{2}(k - 1 + \nu_1^0) \right) \right) \\ &+ \underbrace{\frac{1}{2} \frac{2(1 + k)}{2 + 3k + 3k^2}}_{\theta_2=0} \underbrace{(\rho_1 + (1 - \rho_1)\nu_1^0)}_{\text{prob}[a_1=0|\theta_1=0]} \frac{1}{1 + k} \left(k^2 + k - \frac{1}{2} + \frac{1}{2}\rho_2(0,0) \right) \\ &+ \underbrace{\frac{1}{2} \frac{2(1 + k)}{2 + 3k + 3k^2}}_{\theta_2=1} \rho_1 k \end{aligned}$$

Which is increasing in ν_1^0 for $k > \frac{1}{4}$. Hence the maximum ν_1^0 satisfying equilibrium

conditions is the principal optimal, which corresponds to $\nu_1^0 = \frac{2k\rho_1}{(1 - \rho_1)(1 - 2k)}$ and $\rho_2(0,0) = 1 - 2k$. V_1^P becomes:

$$V_1^P = \frac{k(1 + 12k^3 - 2\rho_1 - 2k^2(4 + \rho_1) - k(1 - 2\rho_1))}{2(-2 + k + 3k^2 + 6k^3)} \quad (\text{E.1})$$

Case 1.2. $\mu_1 = 0$.

Following Proposition 3.2, $\mu_1 = 0$ implies $\nu_1^1 = 0$, i.e., the agents play pooling strategy. As $\rho_1 < 1 - 2k$ and $\rho_2(1, 1) = \rho_1$, $\lambda_2(1, 1) = 0$. δ_1 value is same with Case 1.1 as sequential rationality constraint does not change.

The expected continuation payoff of the principal:

$$\begin{aligned} V_1^P = & \frac{k + 3k^2}{2 + 3k + 3k^2} \left(\rho_1 k + (1 - \rho_1) \left(\frac{1}{2}k + \frac{1}{2}(k - 1 + \nu_1^0) \right) \right) \\ & + \underbrace{\frac{1}{2}}_{\theta_2=0} \underbrace{\frac{2(1+k)}{2 + 3k + 3k^2} (\rho_1 + (1 - \rho_1)\nu_1^0)}_{\text{prob}[a_1=0|\theta_1=0]} \frac{1}{1+k} \left(k^2 + k - \frac{1}{2} + \frac{1}{2}\rho_2(0, 0) \right) \end{aligned}$$

Which is increasing in ν_1^0 for $k > \frac{1}{4}$. Hence the maximum ν_1^0 satisfying equilibrium conditions is the principal optimal, which corresponds to $\nu_1^0 = \frac{2k\rho_1}{(1 - \rho_1)(1 - 2k)}$ and $\rho_2(0, 0) = 1 - 2k$. V_1^P becomes:

$$V_1^P = \frac{k(1 - 8k^2 + 12k^3 - \rho_1 - k - 5k\rho_1)}{2(-2 + k + 3k^2 + 6k^3)} \quad (\text{E.2})$$

Case 2. $\rho_2(0, 0) < 1 - 2k$.

Case 2.2. $\mu_1 > 0$.

The sequential rationality constraint of the bad agent implies:

$$\begin{aligned} (1 - \delta_1)k &= (1 - \delta_1)(k - 1) + \delta_1 \left(\frac{k^2}{1 + k} \right) \\ \Leftrightarrow \delta_1 &= \frac{1 + k}{1 + k + k^2} \end{aligned}$$

For given δ_1 , $\mu_1 = 1$ as $V_2^g > V_2^b$. The expected continuation payoff of the principal:

$$\begin{aligned}
 V_1^P = & \frac{k^2}{1+k+k^2} \left(\rho_1 \left(k - \frac{1}{2} \right) + (1-\rho_1) \left(\frac{1}{2} (k - \nu_1^1) + \frac{1}{2} (k - 1 + \nu_1^0) \right) \right) \\
 & + \underbrace{\frac{1}{2}}_{\theta_1=0} \frac{1+k}{1+k+k^2} \underbrace{(\rho_1 + (1-\rho_1)\nu_1^0)}_{\text{prob}[a_1=0|\theta_1=0]} \frac{k}{(1+k)(1-2k)} \left(2k - 2k^2 - \frac{1}{2} + \frac{1}{2}\rho_2(0,0) \right) \\
 & + \underbrace{\frac{1}{2}}_{\theta_1=1} \frac{1+k}{1+k+k^2} \underbrace{(\rho_1 + (1-\rho_1)\nu_1^1)}_{\text{prob}[a_1=0|\theta_1=1]} \frac{k}{(1+k)(1-2k)} \left(2k - 2k^2 - \frac{1}{2} + \frac{1}{2}\rho_2(0,1) \right)
 \end{aligned}$$

V_1^P is decreasing in ν_1^1 , hence the minimum ν_1^1 satisfying equilibrium conditions is the principal optimal, which corresponds to $\nu_1^1 = \frac{2k\rho_1}{(1-\rho_1)(1-2k)}$ and $\rho_2(0,1) = 1 - 2k$. V_1^P is increasing in ν_1^0 for $k > \frac{1}{4}$. Hence the maximum ν_1^0 satisfying equilibrium conditions is the principal optimal, which corresponds to $\nu_1^0 = \frac{4k\rho_1(1-k)}{(1-\rho_1)(1-2k)^2}$ and $\rho_2(0,0) = (1-2k)^2$. V_1^P becomes:

$$V_1^P = \frac{k^2(-1+6k-12k^2+8k^3+\rho_1)}{2(1-2k)^2(1+k+k^2)} \quad (\text{E.3})$$

Case 2.2.

Following Proposition 3.2, $\mu_1 = 0$ implies $\nu_1^1 = 0$, i.e., the agents play pooling strategy. As $\rho_1 < 1 - 2k$ and $\rho_2(1,1) = \rho_1$, $\lambda_2(1,1) = 0$. δ_1 value is same with Case 2.1 as sequential rationality constraint does not change. The expected continuation payoff of the principal:

$$\begin{aligned}
 V_1^P = & \frac{k^2}{1+k+k^2} \left(\rho_1 k + (1-\rho_1) \left(\frac{1}{2}k + \frac{1}{2}(k - 1 + \nu_1^0) \right) \right) \\
 & + \underbrace{\frac{1}{2}}_{\theta_1=0} \frac{1+k}{1+k+k^2} \underbrace{(\rho_1 + (1-\rho_1)\nu_1^0)}_{\text{prob}[a_1=0|\theta_1=0]} \frac{k}{(1+k)(1-2k)} \left(2k - 2k^2 - \frac{1}{2} + \frac{1}{2}\rho_2(0,0) \right)
 \end{aligned}$$

V_1^P is increasing in ν_1^0 for $k > \frac{1}{4}$. Hence the maximum ν_1^0 satisfying equilibrium conditions is the principal optimal, which corresponds to $\nu_1^0 = \frac{4k\rho_1(1-k)}{(1-\rho_1)(1-2k)^2}$

and $\rho_2(0, 0) = (1 - 2k)^2$. V_1^P becomes:

$$V_1^P = \frac{k^2(-1 + 6k - 12k^2 + 8k^3 + \rho_1)}{2(1 - 2k)^2(1 + k + k^2)} \quad (\text{E.4})$$

Realize that E.3 = E.4, the principal obtains same payoff in Case 2.1. and Case 2.2. Comparison shows that for the defined range of ρ_1 , $V_1^P(E.3) = V_1^P(E.4) > V_1^P(E.1) > V_1^P(E.2)$. The principal gets a higher payoff in Case 2. Hence the principal sets $\delta_1 = \frac{1+k}{1+k+k^2}$. In words, he prefers a more gradual evolution in the reputation level. As the principal is indifferent between babbling and separating behaviour given $\theta_1 = 1$, we select Case 2.2. as our equilibrium candidate as explained in the text. $\square$

Proof. Proposition 3.3

Note that there is no reputational incentive in any equilibria. Suppose that $\nu_i^{\theta_i} > 0$ for some θ_i . The maximum continuation payoff the bad agent can get is $\frac{1}{3}(-\frac{2}{3}) + \frac{1}{3}\frac{1}{3} + \frac{1}{3}\frac{1}{3} = 0$. The agent deviates to play action 1 and get $\frac{1}{3}\frac{1}{3}$. The same reasoning also applies to the good agent. Hence, $\rho_2(0, 0) = 1$ and $\rho_2(1, 1) = \rho_1$. Thus, $\lambda_2(0, 0) = 1$ and $\lambda_2(1, 1)$ depends on V_2^P . For the lowest ρ_2 , $V_2^P = \frac{1}{9}(\rho_2 + \frac{1}{2}\rho_2 - \frac{1}{2}(1 - \rho_2)) = \frac{1}{9}(2\rho_2 - \frac{1}{2})$. As $\rho_2(1, 1) = \rho_1 = 0.1$, $\lambda_2(1, 1) = 0$. The agent is hired in the second period only after $(0, 0)$. Thus the principal's expected continuation payoff:

$$V_1^P = \frac{1}{9} \left(\rho_1 - \frac{1}{2}(1 - \rho_1) + \frac{1}{2}\rho_1 2 \right) = \frac{1}{9} \left(\frac{5}{2}\rho_1 - \frac{1}{2} \right)$$

which is equal to -0.028 for $\rho_1 = 0.1$. $\square$

Proof. Proposition 3.4

There is no reputation incentive in the second period. The reasoning is that playing action 0 can lead a maximum continuation payoff $\frac{1}{21}(-4 + 4) = 0$, the agent will deviate to lay action 1 and get $\frac{2}{21}$. Thus, $\nu_2^0 \& \nu_2^1 \& \mu_2 = 0$.

First we will show that $\lambda_3(1,1) \in (0,1)$ in any equilibrium. Suppose that $\lambda_3(1,1) = 0$. Then, $V_2^b = \frac{2}{21}$. Sequential rationality constraint of the bad agent implies $\nu_1^0 \& \nu_1^1 = 0$. But then, $\rho_2(\theta_1, 0) = 1$, hence $\lambda_3(1,1) = 1$, contradiction.

Suppose that $\lambda_3(1,1) = 1$. Then, $V_2^b = \frac{1}{21}(2+4)$. Sequential rationality constraint of the bad agent implies $\nu_1^0 \& \nu_1^1 = 1$. But then, $\rho_2(0, \theta_1) = \rho_1$ as $\nu_1^1 = 1$ implies $\mu_1 = 1$ by Proposition 3.1. As $\rho_1 = 0.1$, $\rho_3(1,1) = \rho_2 = \rho_1 = 0.1$, then $\lambda_3(1,1) = 0$, contradiction.

For $\lambda_3(1,1) \in (0,1)$, $\rho_3(1,1) = \rho_2 = \frac{1}{3}$. For $\rho_2(0,0) = \frac{1}{3}$, $\frac{\rho_1}{\rho_1 + (1-\rho_1)\nu_1^0} = \frac{1}{3}$. Hence, $\nu_1^0 = \frac{2\rho_1}{1-\rho_1} = \frac{2}{9}$.

There are two possibilities when $\theta_1 = 1$. The principal may hire or not hire the agent after $(1,0)$. Note that the agent will never be hired after $(1,1)$ because $\rho_2(1,1) = \rho_1$. If the principal is hiring the agent after $(1,0)$, then $\nu_1^1 > 0$ and $\mu_1 = 1$ by Proposition 3.1. Moreover, the principal may never hire after observing $\theta_1 = 1$ irrespective of the action. The agent's best response is $\mu_1 = 0$ and $\nu_1^1 = 0$.

Case 1. $\lambda_2(1,0) = 1$

As $\mu_1 = 1$, $\nu_1^1 = \nu_1^0 = \frac{2\rho_1}{1-\rho_1} = \frac{2}{9}$.

$$V_2^P = \frac{1}{21} \left(2\rho_2 + (1-\rho_2)(-1) + \frac{1}{2}\rho_2 4 \right) = \frac{1}{21} (5\rho_2 - 1) = \frac{1}{21} \frac{2}{3}$$

$$V_1^P = \frac{1}{21} \left(-\frac{1}{2} + \left(\rho_1 + (1-\rho_1) \frac{2\rho_1}{(1-\rho_1)} \right) \frac{2}{3} \right) = \frac{1}{21} \left(2\rho_1 - \frac{1}{2} \right) = -\frac{1}{70}$$

Case 2. $\lambda_2(1,0) = 0$

Then, $\mu_1 \& \nu_1^1 = 0$, $\nu_1^0 = \frac{2\rho_1}{1-\rho_1} = \frac{2}{9}$

$$\begin{aligned} V_1^P &= \frac{1}{21} \left\{ \rho_1 + \frac{1}{2}(1-\rho_1)1 + \frac{1}{2}(1-\rho_1)(\nu_1^0 1 - (1-\nu_1^0)2) \right. \\ &\quad \left. + \left(\rho_1 + (1-\rho_1) \frac{2\rho_1}{(1-\rho_1)} \right) \frac{2}{3} \right\} = \frac{1}{21} \left(\frac{11}{2}\rho_1 - \frac{1}{2} \right) = \frac{1}{21} \frac{5}{100} \simeq 0.0024 \end{aligned}$$

As the second case leads to a higher expected continuation payoff, the principal will not hire after observing $\theta_1 = 1$. $\square$

Proof. Proposition 3.5

First we will show that $\lambda_2(1, 1) \& \lambda_3(1, 1) = 0$ in any equilibrium.

1. Suppose that $\lambda_2(1, 1) = 0$ and $\lambda_3(1, 1) > 0$. Then, $V_2^b > \frac{1}{39}3$, hence $\nu_1^{\theta_1} = 1$ for at least one θ_1 . But then $\rho_2(\theta_1, 0) = \rho_1 = \rho_3(1, 1)$, contradicting $\lambda_3(1, 1) > 0$.
2. Suppose that $\lambda_2(1, 1) > 0$ and $\lambda_3(1, 1) > 0$. Note that $\rho_2(1, 1) \leq \rho_1$, hence $\lambda_3(1, 1) = 0$ after $(\theta_1, a_1) = (1, 1)$. Suppose $\lambda_3(1, 1) > 0$ after after $(\theta_1, a_1) = (\theta_1, 0)$ for any θ_1 . Then, $\nu_1^{\theta_1} = 1$ and $\rho_2(\theta_1, 0) = \rho_1$, contradicting $\lambda_3(1, 1) > 0$.
3. Suppose that $\lambda_2(1, 1) > 0$ and $\lambda_3(1, 1) = 0$. $\lambda_3(1, 1) = 0$ implies $V_2^b = \frac{1}{39}3$. If $\lambda_2(1, 1) > 0$, then $\nu_1^1 = 0$. But then, $\mu_1 = 1$ and $\rho_2(1, 0) = 1$. The bad agent will deviate to action 0. Hence there is no such equilibrium.

Following Proposition 3.2, the principal is indifferent between pooling and separating behavior given $\theta_i = 1$ and we choose babbling behavior for the definition of equilibrium.

$$39V_2^P = \left(3\rho_2 + \frac{1}{2}(1 - \rho_2)3 + \frac{1}{2}(1 - \rho_2)(3\nu_2^0 - 6(1 - \nu_2^0)) \right) + \frac{1}{2}(\rho_2 + (1 - \rho_2)\nu_2^0) \left(\frac{27}{2}\rho_3(0, 0) - \frac{9}{2} \right)$$

$$\frac{\partial 39V_2^P}{\partial \nu_2^0} = \frac{9}{2}(1 - \rho_2) - \frac{9}{4}(1 - \rho_2) > 0 \text{ for } \rho_2 < 1$$

Hence, V_2^P is increasing in ν_2^0 . The principal-optimal ν_2^0 is the highest probability satisfying the equilibrium condition so that $\lambda_3(0, 0) = 1$. In particular, $\rho_3(0, 0) =$

$$\frac{\rho_2}{\rho_2 + (1 - \rho_2)\nu_2^0} = \frac{1}{3} \text{ and } \nu_2^0 = \frac{2\rho_2}{1 - \rho_2}. V_2^P \text{ becomes } \frac{1}{26}(9\rho_2 - 1).$$

$\lambda_3(1, 1) = 0$ and $\lambda_2(0, 0) = 1$ implies that $\frac{1}{9} \leq \rho_2(0, 0) \leq \frac{1}{3}$.

$$39V_1^P = \left(-\frac{1}{2}\rho_1 + \frac{1}{2}(1 - \rho_1) + \frac{1}{2}(1 - \rho_1)(\nu_1^0 - 2(1 - \nu_1^0)) \right) + \frac{1}{2}(\rho_1 + (1 - \rho_1)\nu_1^0) \left(\frac{3}{2}(9\rho_2(0, 0) - 1) \right)$$

$$\frac{\partial 39V_1^P}{\partial \nu_1^0} = \frac{3}{2}(1 - \rho_1) - \frac{3}{4}(1 - \rho_1) > 0$$

Hence, V_1^P is increasing in ν_1^0 . The principal-optimal ν_1^0 is the highest probability satisfying the equilibrium condition. In particular, $\rho_2(0, 0) = \frac{\rho_1}{\rho_1 + (1 - \rho_1)\nu_1^0} = \frac{1}{9}$ and $\nu_1^0 = \frac{8\rho_1}{1-\rho_1} = \frac{8}{9}$. V_1^P becomes:

$$V_1^P = \frac{1}{39} \left(\frac{27}{2}\rho_1 - \frac{1}{2} \right) = \frac{1.7}{78} \simeq 0.0218$$

□

Proof. Proposition 3.6

We will first show that $\lambda_2(\theta_1, 0) < 1$ for $\theta_1 \in \{0, 1\}$. Suppose for contradiction that $\lambda_2(\theta_1, 0) = 1$ for some θ_1 . Then $\nu_1^{\theta_1} = 1$ which implies $\rho_2(\theta_1, 0) = \rho_1$. After $(\theta_1, 0)$, if $\lambda_3(\theta_2, 0) = 1$ for some θ_2 , then $\nu_2^{\theta_2} = 1$ and $\rho_3(\theta_2, 0) = \rho_2(\theta_1, 0) = \rho_1 < \frac{1}{3}$, contradicting $\lambda_3(\theta_2, 0) = 1$. Hence, if $\lambda_2(\theta_1, 0) = 1$ for some θ_1 , then $\lambda_3(\theta_2, 0) \in (0, 1)$ for both θ_2 after observing $(\theta_1, 0)$. Then:

$$\begin{aligned} 63V_2^P &= \frac{1}{2}\rho_2 4 + \frac{1}{2}\rho_2 (-8\mu_2 + 4(1 - \mu_2)) + \frac{1}{2}(1 - \rho_2) (-8\nu_2^1 + 4(1 - \nu_2^1)) \\ &\quad + \frac{1}{2}(1 - \rho_2) (4\nu_2^0 - 8(1 - \nu_2^0)) \end{aligned}$$

$\lambda_3(\theta_2, 0) \in (0, 1)$ implies $\rho_3(\theta_2, 0) = \frac{1}{3}$. Hence, $\nu_2^0 = \frac{2\rho_2}{1-\rho_2}$ and $\nu_2^1 = \frac{2\mu_2\rho_2}{1-\rho_2}$ (or $\nu_2^1 = \mu_2 = 0$ as discussed earlier).

$$63V_2^P = 18\rho_2 - 18\rho_2\mu_2 - 2$$

Which is negative for $\rho_2 < \frac{1}{9}$. Thus, $\rho_1 = 0.1$ implies $\lambda_2(\theta_1, 0) < 1$ for both θ_1 . The lowest ρ_2 for which the agent is hired in the second period is $\rho_2 = \frac{1}{9}$ where $\mu_2 = 0$ and $\lambda_3(1, a_2) = 0$ for both a_2 .

Case 1 $\mu_1 > 0$, i.e., separating given $\theta_1 = 1$:

$\lambda_2(\theta_1, 0) \in (0, 1)$ implies $\rho_2(\theta_1, 0) = \frac{1}{9}$ and $\nu_1^{\theta_1} = \frac{8\rho_1}{1-\rho_1} = \frac{8}{9}$ for both θ_1 . It in turn implies $\mu_1 = 1$ by Proposition 3.1. $\lambda_3(0, 0) \in (0, 1)$ implies $\nu_2^0 = \frac{2\rho_2}{1-\rho_2} = \frac{1}{4}$. $\lambda_{i+1}(\theta_i, 0)$ probabilities follow from the sequential rationality constraints of the bad agent.

$\lambda_2(\theta_1, 0) \in (0, 1)$ implies $V_2^P = 0$. Hence, $V_1^P = -\frac{1}{126} \simeq -0.00794$

Case 2 $\mu_1 = \nu_1^1 = 0$, i.e., pooling given $\theta_1 = 1$:

$\lambda_2(\theta_1, 0) \in (0, 1)$ implies $\rho_2(0, 0) = \frac{1}{9}$ and $\nu_1^0 = \frac{8\rho_1}{1-\rho_1} = \frac{8}{9}$. $\lambda_3(0, 0) \in (0, 1)$ implies $\nu_2^0 = \frac{2\rho_2}{1-\rho_2} = \frac{1}{4}$. $\lambda_{i+1}(\theta_i, 0)$ probabilities follow from the sequential rationality constraints of the bad agent.

$\lambda_2(\theta_1, 0) \in (0, 1)$ implies $V_2^P = 0$. Hence,

$$V_1^P = \frac{1}{21} \left(\frac{1}{3}\rho_1 + (1 - \rho_1) \frac{1}{2} \frac{1}{3} + \frac{1}{2} \left(\frac{8}{9} \frac{1}{3} + \frac{1}{9} \frac{(-2)}{3} \right) \right) \simeq 0.013$$

Clearly, the principal prefers the babbling behavior. □