

Do Kings Keep their Promises? A Principal-Agent Problem with Limited Commitment

Andong Yan

Abstract

This research focuses on the principal-agent problem when both principal and agent are constrained only by limited commitment. We parameterize the commitment environment by two factors, the probability of potential contract breach and the cost of contract default. In contrast to the conventional wisdom that lack of commitment (high chance and low cost for contract default) would harm parties' benefits in the contracting, we find that the principal could obtain positive marginal benefits with a higher probability of contract breach, particularly when the costs associated with violating the contract are relatively low. The driving forces behind this unexpected result are that the potential threat of contract breach could behave as a screening tool to separate the agent in their reporting strategy, which leads to a more efficient payment scheme for the principal in the equilibrium.

1 Introduction

The principal-agent problem under limited commitment, involving a powerful principal (such as an autocrat, the central government or employer) delegating tasks to an agent (such as a bureaucrat or employee) whose interests may not align with those of the principal. This misalignment often leads to inefficiencies and requires the design of institution to ensure that the agent acts in the principal’s best interest. The complexities were exacerbated by limited commitment, where neither party can fully commit to future actions, leading to potential conflicts and suboptimal outcomes.

This research focuses on the principal-agent problem when both principal and agent are constrained by limited commitment. We parameterize the commitment environment by two factors: the probability of potential contract default and the cost of contract default. Contrary to the conventional wisdom that a lack of commitment (high chance and low cost for contract default) would harm the parties’ benefits in contracting, we find that the principal could obtain positive marginal benefits with a higher probability of contract breach, particularly when the costs associated with violating the contract are relatively low. The driving force behind this unexpected result is that the potential threat of contract breach can act as a screening tool, influencing the agent’s reporting strategy and leading to a more efficient payment scheme for the principal in equilibrium.

Our study makes contributions to three main branches of literature. First, our work is related to the limited commitment principal-agent problem, particularly in political and economic contexts. Previous studies (Greif, 1993; Myerson, 2015; Acemoglu and Robinson, 2000; Acemoglu, 2003) have highlighted how limited commitment shapes economic institutions and governance. Our research builds on this foundation by parameterizing the commitment environment using the probability of the agent breaking the promise and the cost associated with a costly audit. We find that the principal can obtain positive marginal benefits from a higher probability of breaking the promise given to the agent, particularly when the costs of a costly audit are relatively low. The potential threat of breaking the promise acts as a screening tool, influencing the agent’s reporting strategy and leading to a more efficient payment scheme for the principal. The parametric setting also reveals the dynamics of the principal and the agent’s benefits as the factors of commitment problem vary.

The political economy of autocratic regimes presents unique challenges in governance and administrative efficiency. Olson (1993), Tullock (1987), Egorov and Sonin (2011) and many other seminal works have discussed the dynamics of dictatorship and the difficulties autocrats face in maintaining power and extracting resources. Our research contributes to this literature by analyz-

ing the principal-agent problem that mirrors the context of autocratic governance, where limited commitment and costly verification play crucial roles. We demonstrate that autocratic rulers can use the threat of contract breach and the strategic allocation of verification resources to create efficient incentive structures, even in environments where credible commitment is lacking. However, when the commitment problem is extremely serious, the autocrat has to play the strategy that does not involve any verification to prevent paying for the extra premium that is required to incentivize the agent. This result is similar to the work by Ma and Rubin (2019), who explore the paradox of power in Imperial China, where the lack of credible commitment leads to a low wage-low tax equilibrium. This approach provides new insights that under what environment and mechanism that autocrats can maintain control and ensure administrative efficiency.

The role of costly verification in contract design, delegation problem and mechanism design has been extensively studied. Literature (Townsend, 1979; Halac and Yared, 2020; Ben-Porath, Dekel, and Lipman, 2014) have explored how verification costs shape the structure of optimal contracts. In our model, we incorporate costly verification into a principal-agent with limited commitment framework, examining how the threat of breaking the promise and the associated costly audits influence the principal’s strategy. Our findings suggest that the potential for breaking the promise, coupled with low verification costs, can improve the principal’s ability to design efficient contracts by acting as a screening mechanism that incentivizes the agent to report accurately.

The rest of the chapter is organized as below. Section 2 introduces the model setup. Section 3 displays the strategy of the principal and the agent. Section 4 discusses the equilibrium results. Section 5 presents the numerical example. Section 6 concludes.

2 Model Setup

2.1 Preferences and Setup

We consider a one-period principal-agent problem with limited commitment. A principal owns a project which could potentially generates a income of τ conditional on its success, and nothing if it fails. Whether project is successful is solely determined by a binary status of effort: if there is an input of effort, the project will be a success, otherwise, it fails. The cost of effort to the principal is significantly higher than the income, while the principal could hire an agent who is skilled to the project to work on the project.

The cost of the project for the agent c is drawn from a distribution determined by the state. The agent may face one of a finite number of possible states, $\theta \in \Theta = \{\theta_1, \dots, \theta_n\}$, and we denote the distribution for the cost of the project derived from state θ by $F(\cdot, \theta)$ and its density form $f(\cdot, \theta)$. We assume $F(\cdot, \cdot)$ has common support over $[\underline{c}, \bar{c}]$ for all $\theta \in \Theta$, and $\bar{c} \leq \tau$ so that agent is always efficient in implementing the project. We further assume $F(\cdot, \cdot)$ follows the Monotone Likelihood Ratio Property (MLRP) such that for any $k < j$, $x < y$, we have

$$\frac{f(x, \theta_j)}{f(x, \theta_k)} \leq \frac{f(y, \theta_j)}{f(y, \theta_k)}$$

so that a "higher" state θ_j represents a higher chance to draw a higher project cost than a "lower" state θ_k . The prior distribution for the states is a common knowledge, denote it as $\pi = (\pi_1, \dots, \pi_n) \in \mathcal{P}(\Theta)$, *s.t.* $\sum_{i=1}^n \pi_i = 1$, where $\mathcal{P}(\Theta)$ represents the set of all probability distributions on Θ .

The principal hires the agent and promises a reward $r \in \mathbb{R}^+$ after the project turns out to be a success. Thus the principal will have a payoff of $\tau - r$ conditional on a success, and 0 if project fails. The agent will have a payoff of $r - c$ conditional on a success, and $-c$ if project fails. Thus, we are assuming both the principal and the agent are risk neutral. The agent is also endowed with an outside option that could generate a small but position payoff $\epsilon > 0$.

We assume that the agent could not commit to the effort input which is unobserved by the principal, so the effort input is not a contractible term. And the principal could not fully commit to the payment of the reward after the project outcome is realized either, and we assume that the principal could break the contract at his or her will by paying a "audit" cost κ^* . We consider this as an "auditing" attempt to extract any remaining surplus from the agent, that is $\tau - c$. Moreover, we have κ^* is drawn from a distribution with δ probability to be a finite value κ and $1 - \delta$ probability to be $+\infty$. In other words, there is an exogenous chance of δ that the principal finds breaking the contract is feasible. The audit decision is denoted by $q \in \{0, 1\}$ where 0 stands for "no auditing" and 1 represents "auditing".

2.2 Timing and Information

In this section, we consider the sequence of the game and the information set for the principal and the agent at every possible decision node.

First, there is a contract stage. The agent observes the state $\theta_i \in \Theta$ which is a private information. The agent makes a report $\tilde{\theta} \in \Theta$ to the principal, and then the principal promises a reward r to the agent, which will be realized only if the project delivers a success.

Next, in the action stage, the nature draws the cost of the project c from the distribution $F(c, \theta_i)$, and the cost c is directly revealed to the agent. The agent then decides whether to input an effort into the project. The project outcome is realized according to the agent's decision. It is obvious that the agent will only input an effort if $r \geq c + \epsilon$. It is noteworthy that there exists two levels of information asymmetry in the contract stage and the action stage across the principal and the agent. In the first stage, the agent is better informed by having a signal of θ_i to better infer the distribution of c , while in the second stage, c is directly revealed to the agent. We justify this setting by allowing the agent to learn about the project cost after the project was initiated. And then the agent makes the action decision by comparing the known cost and the potential reward.

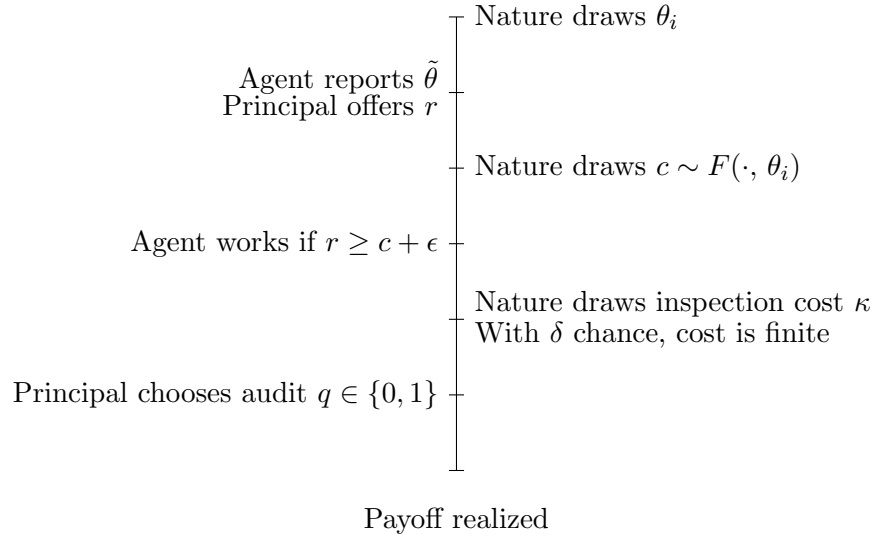


Figure 1: Timeline of the game.

Third, an audit stage. The project outcome is revealed to the principal. Conditional on the success of the project, with δ probability, the audit decision is nontrivial to the principal when the cost of auditing is finite. Then the principal could make a decision to audit or not, q , knowing that the project is successful.

Finally, the payoffs are realized for the principal and the agent. Given a successful project, the principal will have a payoff of $\tau - r$ if there is no auditing, and $\tau - c - \kappa$ if an audit happened; the agent will have a payoff of $r - c$ without being audited, and 0 if audited. If the project fails, both the principal and the agent will have 0 payoff.

Figure 1 summarizes the stages and presents the timeline of the game.

3 Strategy

In the game, the agent makes two decisions, the reporting scheme in the contract stage and the effort input in the action stage. We have stated that the agent will choose to invest efforts only if $r \geq c + \epsilon$, so the action decision is trivial. We only focus on the strategy of the agent's reporting. For a agent who observes the state θ_i , we denote his or her reporting strategy as $\mu_i : \Theta \rightarrow \mathcal{P}(\Theta)$. In other words, the agent could choose a mixed strategy by randomly reporting among possible states $\tilde{\theta}$. Denote the agent's strategy $\{\mu_i\}_{i=1,\dots,n}$ by μ . We further assume that the agent chooses a symmetric strategy: for $i \neq j$, if $\text{supp}(\mu_i) = \text{supp}(\mu_j)$, then we have $\mu_i(\tilde{\theta}) = \mu_j(\tilde{\theta})$ for all $\tilde{\theta}$. That is, if the agent facing different states find a same set of states could yield equally optimal payoffs, the agent will choose the same mixing scheme over the set of states in the reporting strategy.

The principal makes two decisions, the promised reward in the contract stage and the audit decision in the audit stage. The principal could only rely on the report provided by the agent to make decisions. So the principal's strategy is represented by $\sigma = (r, q) : \Theta \rightarrow \mathbb{R}^+ \times \{0, 1\}$, and the principal updates the belief after observing the reports and a conjectured agent's strategy μ by the Bayes' Rule:

$$\phi_k(\tilde{\theta}) = \frac{\pi_k \mu_k(\tilde{\theta})}{\sum_{j=1}^n \pi_j \mu_j(\tilde{\theta})}$$

Given the principal's strategy σ , the agent's problem after observing state θ_i could be represented as:

$$\begin{aligned} \max_{\mu(\cdot)} U(\mu, r, q; \theta_i) &= \sum_{j=1}^n \mu(\tilde{\theta}_j) (1 - \delta q_j) \int_{r_j \geq c + \epsilon} (r_j - c) f(c, \theta_i) dc \\ &= \sum_{j \in q^{-1}(\{0\})} \mu(\tilde{\theta}_j) \int_{r_j \geq c + \epsilon} (r_j - c) f(c, \theta_i) dc + \sum_{j \in q^{-1}(\{1\})} \mu(\tilde{\theta}_j) (1 - \delta) \int_{r_j \geq c + \epsilon} (r_j - c) f(c, \theta_i) dc \end{aligned}$$

The agent takes the expectation over the net payoff on the range of the cost realization that will ensure a project success, and then weight it by the probability of each state that is mixing in the report scheme. We decompose the agent's value function to explicitly show that the agent will discount the expected payoff by $(1 - \delta)$ for the range of reports that the principal will audit ($q^{-1}(\{1\})$). To simplify the notation, we define the integral term in the agent's value function as

below:

$$I_k(r) = \int_{\underline{c}}^{r-\epsilon} (r-c)f(c, \theta_k) dc$$

Given the agent's strategy, the principal maximizes expected utility with sequential rationality constraints

$$\max_{\{r_j, q_j\}_{j=1, \dots, n}} \sum_{i=1}^n \sum_{j=1}^n \pi_i \mu_i(\tilde{\theta}_j) V(\mu, r_j, q_j; \tilde{\theta}_j)$$

where

$$\begin{aligned} V(\mu, r_j, q_j; \tilde{\theta}_j) &= \sum_{k=1}^n \phi_k(\tilde{\theta}_j) \left[(1 - \delta q_j) \int_{r_j \geq c+\epsilon} (\tau - r_j) f(c, \theta_k) dc + \delta q_j \int_{r_j \geq c+\epsilon} (\tau - \hat{c}_k(r_j) - \kappa) f(c, \theta_k) dc \right] \\ \hat{c}_k(r) &= \int c f_{c \leq r}(c, \theta_k) dc \\ q_j &= \arg \max_q \{ \tau - r_j, \tau - \hat{c}_k(r_j) - \kappa \}. \end{aligned}$$

The principal considers the posterior probability of true state being θ_k given the agent's report $\tilde{\theta}_j$ and forms the expected utility. r_j represents the promised reward receiving report $\tilde{\theta}_j$, q_j represents the proposed audit decision, $\hat{c}_k(r)$ represents the expected cost of the project conditional on a success given a reward of r under state θ_k , and the last equality constraint is the sequential rationality constraint for the principal, which requires the principal to have the correct incentive to make the audit decision consistent to the proposed strategy.

We define q is a cut-off audit strategy if there exists $0 \leq \bar{k} \leq n$ such that for $i \in \{1, \dots, n\}$, we have

$$q_i = \begin{cases} 1, & \text{for } i > \bar{k} \\ 0, & \text{for } i \leq \bar{k} \end{cases}$$

and define r is a type-monotone reward strategy if we have

$$r_1 \leq \dots \leq r_i \leq \dots \leq r_n$$

. For the solution concept, we consider the Perfect Bayesian Equilibrium (PBE) and a strategy profile (σ, μ) is an equilibrium if it solves the agent and the principal's problems stated above. We are particularly interested in the subset of the equilibria which has a type-monotone reward

strategy. This is a natural restriction as a higher state represents a higher chance of a higher cost faced by the agent, thus requires the principal to provide stronger incentives. Note that this restriction is not necessary for the equilibrium to sustain, but only helps us to focus on the set of the equilibria that we are more interested in.

4 Equilibrium Results

Our first result states the form of the agent's optimal response μ to the principal's strategy σ . We first define the following mapping $\hat{r}_k(\cdot) : \mathbb{R}^+ \rightarrow \mathbb{R}^+$ such that

$$I_k(r) = (1 - \delta)I_k(\hat{r}_k) \Rightarrow \hat{r}_k(r) = I_k^{-1} \left(\frac{I_k(r)}{1 - \delta} \right)$$

Since $I_k(\cdot)$ is a continuous and increasing function, the mapping is well defined and we have $\hat{r}_k(r) > r$. Essentially, $\hat{r}_k(\cdot)$ pins down the equality of the two terms in the agent's value function, and maps from any reward promised by the principal with no potential audit, to the desired reward amount to yield same expected payoff under the scenario that an audit is proposed. Then we could have the following proposition regarding the agent's decision rule:

Proposition 1. *For agent type θ_i , facing a type-monotone reward strategy and a non-trivial audit strategy ($q_j = 1$ for at least one reported state $\tilde{\theta}_j$) by the principal, the agent will choose to report:*

$$\begin{cases} \mu \left(\left\{ \tilde{\theta}_i \mid r_i \in \max\{r_j \mid j \in q^{-1}(\{0\})\} \right\} \right) = 1 & \text{if } r_{k_1} < \hat{r}_k(r_{k_0}) \\ \mu \left(\left\{ \tilde{\theta}_i \mid r_i \in \max\{r_j \mid j \in q^{-1}(\{1\})\} \right\} \right) = 1 & \text{if } r_{k_1} > \hat{r}_k(r_{k_0}) \\ \mu \left(\left\{ \tilde{\theta}_i \mid r_i \in \left\{ \max\{r_j \mid j \in q^{-1}(\{0\})\} \cup \max\{r_j \mid j \in q^{-1}(\{1\})\} \right\} \right\} \right) = 1 & \text{if } r_{k_1} = \hat{r}_k(r_{k_0}) \end{cases}$$

where $k_0 = \arg \max \{i \in q^{-1}(\{0\})\}$ and $k_1 = \arg \max \{i \in q^{-1}(\{1\})\}$.

Proof. From the agent's value function, we have $I_k(r)$ an increasing function in r , which implies the agent has a dominant strategy to report at the state that promises the highest reward, conditional on the audit strategy. The audit strategy creates two subsets in the state space, namely $q^{-1}(\{0\})$ and $q^{-1}(\{1\})$, and the candidate of states to report could be only from the highest state in each subset, due to the type-monotone reward strategy we are considering here. Finally, we consider the comparison between the expected payoff of the agent in these two scenarios to pin down the optimal reporting scheme. $\square$

Next, we have the second result to show that we could limit our focus to the equilibrium with a cut-off audit strategy:

Proposition 2. *There exists an equilibrium with a cut-off strategy that could result in the outcome derived from any equilibrium.*

Proof. The arbitrary audit strategy q generates $q^{-1}(\{0\})$ and $q^{-1}(\{1\})$. Following proposition 1, we could find k_0 and k_1 . There are only two scenarios regarding the order of k_0 and k_1 .

Suppose $k_0 < k_1$, then the same equilibrium outcome could be achieved by a cut-off audit strategy by setting $\bar{k} = k_0$, in this case $k_1 = n$. The equilibrium outcome is either the agent reports at $\tilde{\theta}_{k_0}$ and other states with same reward, the principal proposes a reward r_{k_0} and $q = 0$, or the agent reports at $\tilde{\theta}_n$ and other states with same reward, the principal proposes a reward r_n and $q = 1$, which depends on the comparison of r_{k_n} and $\hat{r}_k(r_{k_0})$.

Suppose $k_0 > k_1$, by proposition 1, agents will always choose to report at θ_{k_0} or its equivalent states as we have $\hat{r}_k(r_{k_0}) > r_{k_0} > k_1$. To have this equilibrium outcome, we could set the cut-off strategy with $\bar{k} = n$ so that the principal proposes a contract that rules out the possibility of auditing. This will result in the same equilibrium outcome as the original equilibrium. $\square$

Now combining proposition 1 and 2, we could rewrite the principal's strategy as $(r_0, r_1, \bar{k})$ such that

$$q(\tilde{\theta}_i) = \mathbb{I}(i > \bar{k}); (\tilde{\theta}_i) = \begin{cases} r_0 & \text{for } i \leq \bar{k} \\ r_1 & \text{for } i > \bar{k} \end{cases}$$

and also the principal's optimization problem could be represented as:

$$\max V_k \in \{V_0, V_1, \dots, V_n\}$$

where

$$V_k = \max_{r_0, r_1} \sum_{i=1}^k \pi_i \int_{r_0 \geq c+\epsilon} (\tau - r_0) f(c, \theta_i) dc + \sum_{i=k+1}^n \pi_i \left[(1 - \delta) \int_{r_1 \geq c+\epsilon} (\tau - r_1) f(c, \theta_i) dc + \delta \int_{r_1 \geq c+\epsilon} (\tau - \hat{c}_i(r_1) - \kappa) f(c, \theta_i) dc \right]$$

with following constraints

$$r_1 \in [\hat{r}_{\bar{k}+1}(r_0), \hat{r}_{\bar{k}}(r_0)] \quad (\text{Agent IC})$$

$$\mathbb{E}_{\theta_i} [U(\hat{\mu}, r, q; \theta_i)] \geq \epsilon \quad (\text{Agent IR})$$

$$r_0 \leq \kappa + \sum_{i=1}^{\bar{k}} \pi_i \hat{c}_i(r_0) \quad (\text{Principal IC-No Audit})$$

$$r_1 \geq \kappa + \sum_{i=\bar{k}+1}^n \pi_i \hat{c}_i(r_1) \quad (\text{Principal IC-Audit})$$

In the simplified principle's problem, we are searching over all potential thresholds for the cut-off audit strategy, while imposing constraints that ensure the incentive compatibility for both the principal and the agent, and then maximize the principal's expected payoff by controlling the two reward levels within the constrained ranges. We are imposing that the agent will follow the proposed contract by the principal, in the sense that if the agent's type is below the threshold, the agent to choose to take the reward r_0 , otherwise the agent will take the reward r_1 even it is accompany with a potential audit.

The agent's incentive constraint concerns whether the agent will deviate to the opposite reward scheme, e.g. from $(r_0, q = 0)$ to $(r_1, q = 1)$, or vice versa. We could show that $\hat{r}_{i+1}(r_0) > \hat{r}_i(r_0)$ for every i , thus we only need to consider the two types across the threshold of the boundary, that is $\theta_{\bar{k}}$ and $\theta_{\bar{k}+1}$. We restrict $r_1 \leq \hat{r}_{\bar{k}}(r_0)$ which prevents the agent of type $\theta_{\bar{k}}$ from deviating, and $r_1 \geq \hat{r}_{\bar{k}+1}(r_0)$ ensures the incentive compatibility for the type $\theta_{\bar{k}+1}$.

The agent's individual rationality constraint ensures that the agent has an ex-ante expected payoff more than the outside option. And the last two principal's incentive constraints state that the principal should indeed follow the proposed audit decision in the ranges below and above the cut-off threshold.

There are several remarks:

(1) In the principal's value function V_k given the cutoff audit threshold at k , we are aggregating the prior distributions π_i of each state θ_i to formulate the expectation in the ranges below or above the threshold. This is feasible as the proposed reward is a flat value in these two ranges, so that by the agent's symmetric reporting strategy, any reported state within the range is uninformational and the principal could only rely on the prior distribution.

(2) It is possible that there is no feasible solution to the maximization problem of V_k as the constraints are too tight and resulting in no feasible ranges. However, there is a last resort for the equilibrium to exist at V_n , which represents that the principal's strategy involves no audit. This is always feasible by paying a low flat wage for all reported states.

(3) Agent's incentive constraint might be conflicting when $\hat{r}_{k+1}(r_0) > \hat{r}_k(r_0)$, which implies there are no feasible reward scheme to provide incentives to persuade the agent not to deviate to the other types' reward scheme. While we allow this situation to happen in general, we could show that under some further normality assumption on the cost distribution, we could have $\hat{r}_i(r_0) > \hat{r}_j(r_0)$ for every $i > j$, thus the reward scheme to separate between any two types is feasible.

Now we have our main result:

Proposition 3. *The solution to the principal's simplified problem $(r_0, r_1, \bar{k})$, together with the induced agent's reporting scheme, such that the agent will choose to report randomly over $\{\tilde{\theta}_1, \dots, \tilde{\theta}_{\bar{k}}\}$ if $\theta_i \in \{\theta_1, \dots, \theta_{\bar{k}}\}$; and report randomly over $\{\tilde{\theta}_{\bar{k}+1}, \dots, \tilde{\theta}_n\}$ if $\theta_i \in \{\theta_{\bar{k}+1}, \dots, \theta_n\}$, is a PBE.*

Proof. It could be verified that the belief induced by the agent's strategy is consistent with the principal's belief in the optimization problem. The principal's strategy satisfies the sequential rationality which is imposed in the problem. And the principal's strategy and the agent's strategy are the optimal responses to each other. $\square$

4.1 Discussion

We could compare the result in last section with the benchmark when the principal has no commitment problem in the environment. In other words, the principal could commit to a reward scheme ad-hoc. In this scenario, as the agent has a dominant strategy to always choose to report at the state related to the highest reward regardless of the state, there will be a pooling equilibrium that any report is uninformational to the principal, and the principal has to make the decision based on the prior belief.

On the other hand, the result we have displayed in the limited commitment environment shows that the principal could utilize the potential audit to separate the agents into two groups. There is a tradeoff in this new feature of the equilibrium: on one side, the principal could be better off due to the improved efficiency from the information gain of the screening results; on the other side, the principal has to pay a premium of risk to the higher type group to ensure the incentive

compatibility. Moreover, the principal is still limited by his or her own incentive compatibility constraints.

Overall, whether the principal (and the agent) could benefit from the exercise of the audit is ambiguous. Due to the complexity of the analysis of the non-linear optimization problem, we present a numerical analysis in the next section.

5 Numerical Analysis

In this section, we solve the principal's simplified problem numerically, and discuss the analysis of comparative statics. We consider the state θ_i to be drawn from $\Theta = \{\theta_1, \dots, \theta_{10}\}$, with equally likely probability. The cost distribution follows a family of beta distributions, and as i increases, the density distribution assigns more weight from left to the right. The below figure displays the family of beta distributions that we are considering in this example.

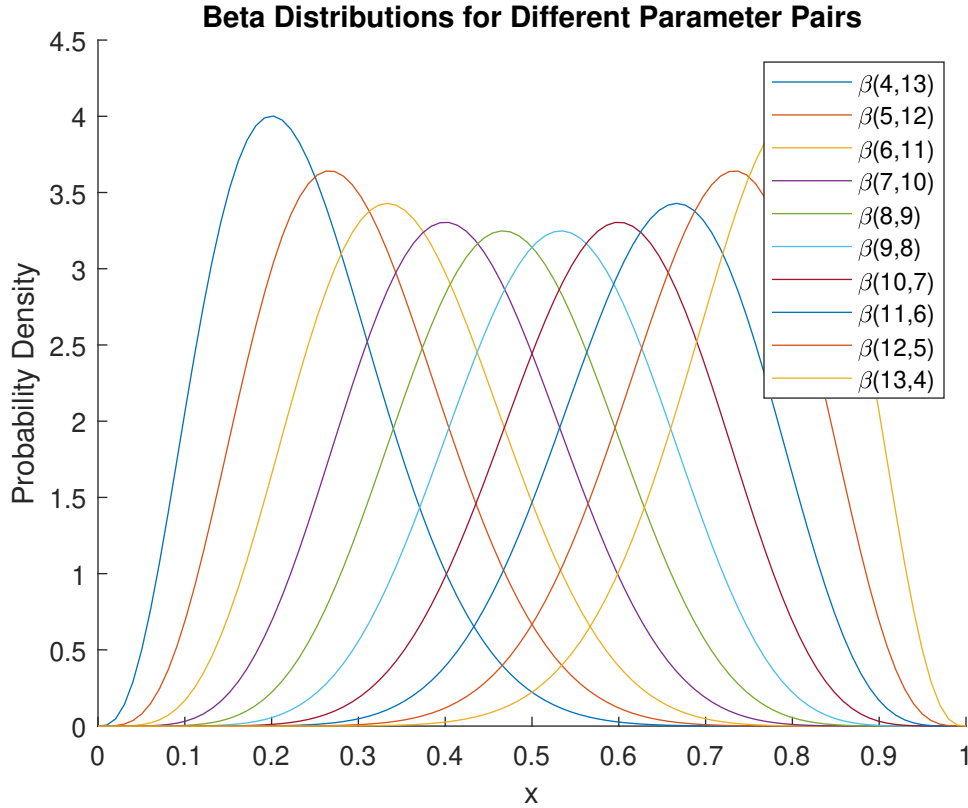


Figure 2: Family of Beta Distribution

We consider $\tau = 1$, $\epsilon = 0.01$, $[\underline{c}, \bar{c}] = [0, 1]$, $\kappa \in \{0.35, 0.4, 0.43\}$ and $\delta \in (0.01, 0.99)$. The

results are presented in figures below. We would like to display the dynamics of the equilibrium in both dimensions of δ and κ . Thus, in each of the panel, we solve for the equilibrium (ex-ante) payoffs for the principal and the agent, as well as the cut-off threshold of the audit strategy, for the range of δ with a step of 0.01. In the three panels, we use value of κ to be 0.35, 0.4 and 0.43, respectively.

One interesting finding is that the principal's payoff exhibits a non-monotone relation as δ increases, when κ is relatively low. This implies when δ level is intermediate, the efficiency improvement effect from the screening of the agents dominates the increased cost of the reward payment. However, the marginal benefits decreases as δ keeps increasing and the principal has to promise more reward to satisfy the agent's incentive constraints.

When δ is close to 1, that is when the commitment problem is very serious, the principal found it not beneficial to use its power anymore, and resulting in the non-audit strategy with threshold $\bar{k} = 10$, as shown in the figure. This implies the payoffs generated from strategies involving auditing is bringing negative marginal benefits so the principal resorts back to the equilibrium without auditing.

We could also find that the agent's utility generally decreases with δ , except at the point where the audit strategy changes. This is natural as when there exists some possibility of auditing, a higher chance of auditing will decrease the agent's payoffs. When the auditing threshold becomes less tight, that is when $\bar{k}$ increases, the agent will benefit from the less scenarios that will be audited. When strategies with non-audit are in place, the agent receives the most payoffs.

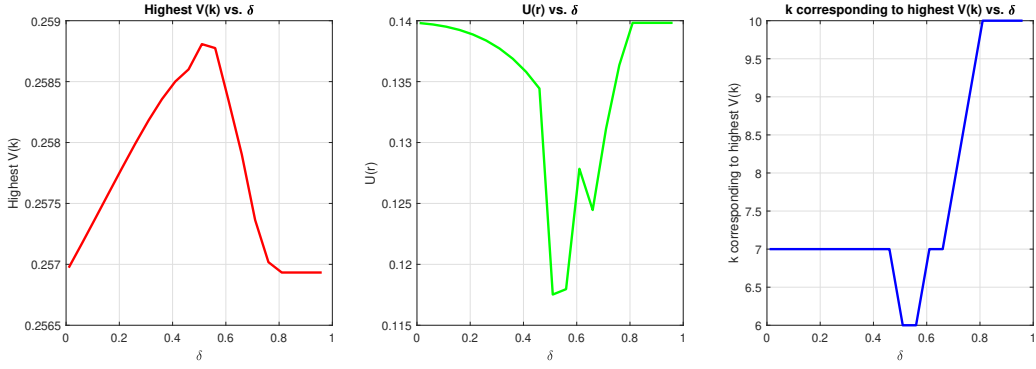


Figure 3: $\kappa = 0.35$, $\tau = 1$, $\epsilon = 0.01$, $[\underline{c}, \bar{c}] = [0, 1]$

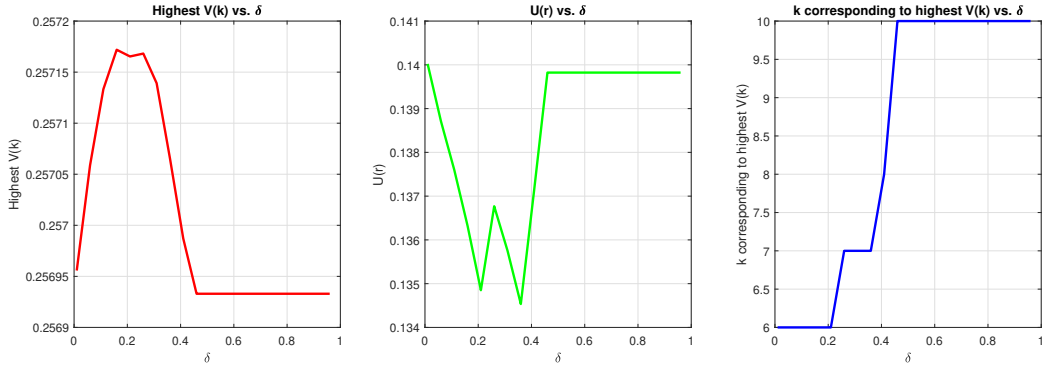


Figure 4: $\kappa = 0.4$, $\tau = 1$, $\epsilon = 0.01$, $[\underline{c}, \bar{c}] = [0, 1]$

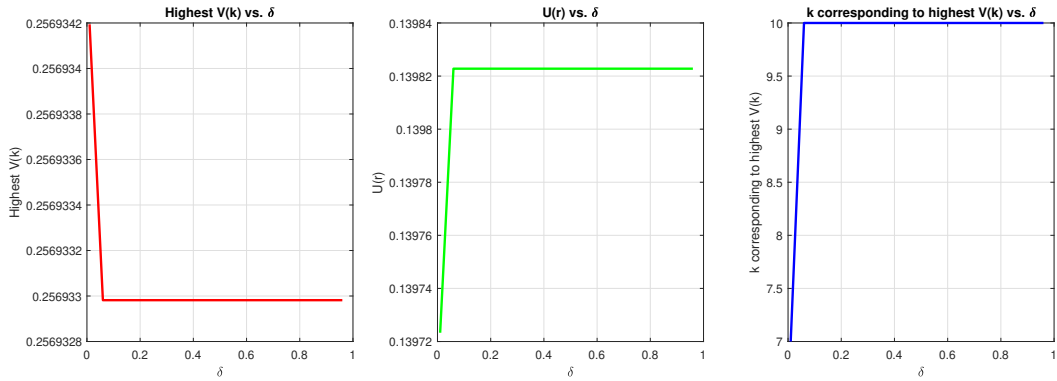


Figure 5: $\kappa = 0.43$, $\tau = 1$, $\epsilon = 0.01$, $[\underline{c}, \bar{c}] = [0, 1]$

6 Conclusion

This research examines the principal-agent problem under limited commitment, focusing on the interaction between a principal and an agent constrained by their inability to fully commit to future actions. By parameterizing the commitment environment through the probability of contract default and the associated costs, we provide a counterexample to the conventional wisdom that limited commitment generally harms contracting parties. Our findings reveal that a higher probability of contract breach, particularly when the costs of violation are low, can provide marginal benefits to the principal because the threat of contract breach acts as a screening tool, enhancing the efficiency of the payment scheme. This study extends to the political economy of autocratic regimes, demonstrating how autocratic rulers can use the threat of contract breach and strategic allocation of verification resources to create efficient incentive structures and maintain administrative efficiency, even in environments where credible commitment is lacking. These insights offer valuable implications for understanding the dynamics of governance and policy implementation, especially in autocratic settings, highlighting the nuanced strategies that can be employed to address the challenges posed by limited commitment.

7 Reference

Halac, M., & Yared, P. (2020). Commitment versus Flexibility with Costly Verification. *Journal of Political Economy*, 128(12), 4523–4573. <https://doi.org/10.1086/710560>

Townsend, R. M. (1979). Optimal contracts and competitive markets with costly state verification. *Journal of Economic Theory*, 21(2), 265–293. [https://doi.org/10.1016/0022-0531\(79\)90031-0](https://doi.org/10.1016/0022-0531(79)90031-0)

Ben-Porath, E., Dekel, E., & Lipman, B. L. (2014). Optimal Allocation with Costly Verification. *American Economic Review*, 104(12), 3779–3813. <https://doi.org/10.1257/aer.104.12.3779>

Erlanson, A., & Kleiner, A. (2020). Costly verification in collective decisions. Unpublished paper. Available at [arXiv:1910.13979](https://arxiv.org/abs/1910.13979).

Greif, A. (1993). Contract enforceability and economic institutions in early trade: The Maghribi Traders’ Coalition. *American Economic Review*, 83(3), 525–548.

Myerson, R. B. (2015). Moral hazard in high office and the dynamics of aristocracy. *Econometrica*, 83(6), 2083–2126. <https://doi.org/10.3982/ecta9737>

Myerson, R. B. (2008). The autocrat’s credibility problem and foundations of the constitutional state. *American Political Science Review*, 102(1), 125–139. <https://doi.org/10.1017/s0003055408080076>

Egorov, G., & Sonin, K. (2011). Dictators and Their Viziers: Endogenizing the Loyalty-Competence Trade-Off. *Journal of the European Economic Association*, 9(5), 903–930.

Acemoglu, D., & Robinson, J. A. (2000). Political losers as a barrier to economic development. *American Economic Review*, 90(2), 126–130. <https://doi.org/10.1257/aer.90.2.126>

Miquel, G. P. I., & Yared, P. (2012). The political economy of indirect control. *Quarterly Journal of Economics*, 127(2), 947–1015. <https://doi.org/10.1093/qje/qjs012>

Qian, Y., & Weingast, B. R. (1997). Federalism as a commitment to preserving market incentives. *Journal of Economic Perspectives*, 11(4), 83–92. <https://doi.org/10.1257/jep.11.4.83>

Olson, Mancur (1993). “Dictatorship, Democracy, and Development.” *American Political Science Review*, 87, 567–576.

Tullock, Gordon (1987). *Autocracy*. Kluwer Academic.

Tilly, C. (1990). *Coercion, Capital, and European States, AD 990-1990*. Wiley-Blackwell.