

Repeated Delegation

Elliot Lipnowski*
University of Chicago

João Ramos†
University of Southern California

August 2017

Abstract

We study an ongoing relationship of delegated decision making. Facing a stream of projects to potentially finance, a principal must rely on an agent to assess the returns of different opportunities; the agent has lower standards, wishing to adopt every project. In equilibrium, the principal allows bad projects in the future to reward fiscal restraint by the agent today, but she cannot commit to reward the agent indefinitely. We fully characterize the equilibrium payoff set (at fixed discounting), showing that Pareto optimal equilibria can be implemented via a two-regime ‘Dynamic Capital Budget’. We show that, rather than backloaded rewards—a prevalent feature of dynamic agency models with greater commitment power—our Pareto optimal equilibria feature an inevitable loss of autonomy for the agent as time progresses. This transition toward conservatism speaks to the life cycle of an organization: as it matures, it generates lower revenue at a higher yield.

Keywords: Delegation, limited commitment, repeated game, capital budgeting

JEL codes: C73, D23, D82, D86, G31

*E-mail: lipnowski@uchicago.edu and joao.ramos@marshall.usc.edu.

†We thank Joyee Deb, David Pearce, and Ennio Stacchetti for valuable insights and constant encouragement. This work has also benefited from discussions with Heski Bar-Isaac, V. Bhaskar, Adam Brandenburger, Ben Brooks, Sylvain Chassang, Peter DeMarzo, Ross Doppelt, Eduardo Faingold, Alex Frankel, Vijay Krishna, Laurent Mathevet, Debraj Ray, Evan Sadler, Tomasz Szadzik, Maher Said, Larry Samuelson, Alex Wolitzky, Chris Woolnough, Sevgi Yuksel, and seminar participants at NYU, Yale, UWaterloo, FGV, USP, McGill, BU, UWO, USC Marshall, UChicago, UToronto, Northwestern, Princeton, Rochester, UPenn, Econ-Con, Midwest Theory, Public Economic Theory, Canadian Economic Theory, the ES winter meetings, SAET, and SITE. We thank Eric Karsten for research assistance.

1 Introduction

Many economic activities are arranged via delegated decision making. In practice, those with the necessary information to make a decision may differ—and, indeed, have different interests—from those with the legal authority to act. Such relationships are often ongoing, consisting of many distinct decisions to be made over time, with the conflict of interest persisting throughout. A state government that funds local infrastructure may be more selective than the local government equipped to evaluate its potential benefits. A university bears the cost of a hired professor, relying on the department to determine candidates' quality. The Department of Defense funds specialized equipment for each of its units, but must rely on those on the ground to assess their need for it. Our focus is on how such repeated delegation should optimally be organized, and on how the relationship evolves over time.

The above relationships, hindered by conflicting interests, also suffer from limited means to align those interests. Beyond the limits of monetary incentives—indeed, the state government would not have the mayor pay out of pocket to build a city park—formal contingent contracting may be difficult for two reasons. First, it may be impractical for the informed party to produce verifiable evidence supporting its recommendations. Second, it might be unrealistic for the controlling party to credibly cede authority in the long-run, i.e. to commit. Even so, the prospect of a future relationship may align the actors' interests: both parties may be flexible concerning their immediate goals, with a view to a healthy relationship.

We study an infinitely repeated game in which a principal (“she”) has authority over a decision to be made in an uncertain world; she relies on an agent (“he”) to assess the state. Each period, the principal must choose whether or not to initiate a project, which may be good (i.e. high enough value to offset its cost) or bad. The principal herself is ignorant of the current project's quality, but the agent can assess it. The players have partially aligned preferences: both prefer a good project to any other outcome, but they disagree on which projects are worth adopting. The principal wishes to fund only good projects, while the agent always prefers to fund. For instance, consider the ongoing relationship between local and state governments. Each year, a county can request state government funds for the construction of a park. The state, taking into account past funding decisions, decides whether or not to fund it. The park would surely benefit the county, but the state must weigh this benefit against the money's opportunity cost. To assess this tradeoff, the state

relies on the county's local expertise. We focus on the case in which the principal needs the agent: the ex-ante expected value of a project is not enough to offset its cost. Thus, if the county were never selective in its proposals, the state would not want to fund every park.

To delegate—to cede control at the ex-ante stage—entails some vulnerability. If our principal wants to make use of the agent's expertise, she must give him the liberty to act. This vulnerability limits the freedom that the agent can expect from the principal in the future. In funding a park, the state government risks wasting taxpayer money. In particular, if the state makes a policy of funding each and every park the county requests, then it risks wasting money on many unneeded parks. Said differently, when the time comes to fund those parks, the state government would rather renege on that implicit agreement. The state government cannot credibly reward a county's fiscal restraint today by promising *carte blanche* in the future.

Our first main result says that the delegation relationship becomes more conservative as time progresses. In the early stages of the relationship, the principal always cedes control to the agent, who in turn adopts every good project that arrives but occasionally adopts bad projects as well. The agent clearly benefits from this freedom, but the principal faces a genuine tradeoff: while many projects are adopted, the average adopted project has a low yield, i.e. value per unit cost. Eventually, the agent's goodwill permanently runs dry. After this happens, the principal rarely delegates; the agent may adopt some good projects if allowed, but will never again squander an opportunity on a bad project.

Thus, combining three sensible ingredients—limited information, limited liability, and limited commitment—yields meaningful, new conclusions for the evolution of the relationship. Indeed, if the agent's cost of exercising restraint were observable, we would expect him to be rewarded over time, as in Ray (2002).¹ Maintaining unobservability, some papers focus on applications for which principal commitment power is appropriate. Guo and Hörner (2015), for instance, obtain a remarkably complete characterization of the optimal contract, even in the case of persistent valuations. Their optimal contract yields divergent long-run outcomes, permanently punishing or rewarding the agent based on early random occurrences—even if the principal would rather abandon the relationship than deliver this reward. By contrast, in situations where the controlling party retains control over time—and therefore cannot pre-commit to unlimited spending—we should expect a transition to a more conservative relationship.

¹Ray (2002) provides a general result capturing this common phenomenon. See also Lazear (1981) and Harris and Holmström (1982).

Our second main result is a complete characterization of the equilibrium payoff set. If the principal faced no credibility concerns—as is the case in Guo and Hörner (2015) and Li, Matouschek, and Powell (2017)—the techniques of Spear and Srivastava (1987) would apply.² Accordingly, those two papers obtain the set of possible payoffs via a dynamic program, with the agent’s promised value as the state variable. This approach is not appropriate for our setting. As the principal would rather renege when called to deliver too many bad projects, there are (endogenous) limits to the value the agent can be promised. Lacking a general method to explicitly compute this set for a repeated game of incomplete information (at fixed discounting), we must combine the recursive methods of Abreu, Pearce, and Stacchetti (1990) with specific features of our stage game to arrive at a tractable problem.

The above characterization yields our third main result, which uncovers the form of the optimal intertemporal delegation rule. In keeping with the intuition (as in, say, Jackson and Sonnenschein (2007)) that linking decisions aligns incentives, delegation is organized via a budgeting rule. The uniquely³ principal-optimal equilibrium, the **Dynamic Capital Budget**, comprises two distinct regimes. At any time, the parties engage in either Capped Budgeting or Controlled Budgeting.

In the **Capped Budget** regime, the principal always delegates, and the agent initiates all good projects that arrive. At the relationship’s outset, the agent has an expense account for projects, indexed by an initial balance and an account balance cap. The balance captures the number of projects that the agent could adopt immediately without consulting the principal. Any time the agent takes a project, his balance declines by 1. While the agent has any funds in his account, the account accrues interest. If the agent takes few enough projects, the account will grow to its cap. At the cap, the agent is still allowed to take projects, but his account grows no larger (even if he waits). Not being rewarded for fiscal restraint, the agent immediately initiates a project, and his balance again declines by 1.

If the agent overspends, a **Controlled Budget** regime begins: the principal only delegates sometimes in this regime, but the agent never adopts a bad project again. This inevitable outcome—indeed, Theorem 1 tells us Capped Budgeting can only be temporary—is a unfortunate one, Pareto inferior to the Capped Budget regime. This dismal future, with the threat it entails, is necessary to sustain optimal delegation in the earlier stages of the

²The principal of Guo and Hörner (2015) has commitment power by assumption, as discussed above. Meanwhile, the principal of Li, Matouschek, and Powell (2017) can credibly promise any value to the agent, as the agent has a means (via effort decisions) to punish her otherwise.

³More precisely, the two-regime structure on path and the exact character of the Capped Budget regime are uniquely required by optimality.

relationship.

Related Literature This paper studies delegated decision making.⁴ A principal must rely on a biased, informed agent to make a decision, but the principal cannot use monetary compensation or other ex-post rewards to incentivize the agent. In deciding how much freedom to give, the principal faces a tradeoff between leveraging the agent’s private information and shielding herself from his conflicting interests. A key insight from this literature is that some ex-post inefficiency may be optimal, providing better incentives to the agent ex-ante.

This insight leads to a natural question when parties interact over time: How should the future of a delegation-based relationship be distorted to provide incentives today? We join a growing literature on dynamic delegation which seeks to answer this question. Most closely related are Guo and Hörner (2015) and Li, Matouschek, and Powell (2017).⁵ Guo and Hörner (2015) study optimal dynamic mechanisms without money in a setting where the principal can commit. The paper first considers a setting in which the agent’s valuation is independent across time. Each period, the uninformed principal decides whether to provide a costly, perishable good to the agent. This setting is essentially the same as our model, but with the key difference that the principal can commit ex-ante to an intertemporal allocation rule. The benchmark model of Li, Matouschek, and Powell (2017) is a repeated game of project selection: each period, a biased agent and an uninformed principal together select a project, at which point they simultaneously choose an implementation effort. This latter effort decision gives their principal effective commitment power, by giving each player a way to unilaterally punish the other.⁶ In each of the above two models, the optimal mechanism generates path dependence in long-run outcomes. The players enter one of two absorbing regimes (each with positive probability): one where the agent suffers, and one where the agent receives his first-best outcome indefinitely. Thus, given principal commitment, early random events have long-lasting consequences for the relationship.

The key contribution of our paper to this literature is to understand dynamic delegation absent commitment power. At a technical level, this poses a challenge. With commitment power, it is well-known (see Spear and Srivastava (1987)) that the optimal contract may be derived as the solution of a dynamic program, using the agent’s promised utility as a state

⁴The seminal work is by Holmström (1984). For more recent work, see Armstrong and Vickers (2010) and Ambrus and Egorov (2017), as well as Frankel (2014) and the thorough review therein.

⁵See Guo (2016), Alonso and Matouschek (2007), and Bird and Frug (2017) as well.

⁶Indeed, the optimal contracting analogue of this model—wherein the principal faces no incentive-compatibility constraints—generates the same path of play as their principal-optimal equilibrium.

variable. This result, employed both by Guo and Hörner (2015) and Li, Matouschek, and Powell (2017), is inapplicable in settings of limited commitment. Instead, we must work directly with the machinery of Abreu, Pearce, and Stacchetti (1990) for repeated games and work to simplify the principal’s incentive constraints; this approach is also taken by Fong and Li (2017) in an employment setting. More than requiring different methods, our limited commitment setting yields substantively different economic predictions for dynamic delegation. Rather than facing divergent outcomes as in these other papers, our players’ relationship is certain to become more conservative. Our principal, unable to credibly reward her agent in the long-run, must balance long-run punishment and medium-term reward to elicit good behavior today.

Our results add to the literature on relationship building under private information. While focusing on a different misalignment of preferences, Möbius (2001) and Hauser and Hopenhayn (2008) look at a model of trading favors, in which a player’s opportunity to do a favor is private.⁷ The form of conflict is qualitatively different there: unlike in our model (wherein adopting good projects benefits everybody), every action that benefits one player harms another in the stage game. Hauser and Hopenhayn (2008) show that the relationship can benefit from varying incentives based on both action and inaction, a feature which reappears in our model. Their simulations indicate that forgiveness is a feature of every Pareto optimal equilibrium. In our model of partially aligned preferences, there are fundamental limits to forgiveness: our players eventually permanently face a Pareto dominated outcome.

The allocation of future responsibility is a common incentivizing device in the dynamic corporate finance literature. In Biais et al. (2010), for instance, the principal commits to investment choices and monetary transfers to the agent, who privately acts to reduce the chance of large losses for the firm.⁸ While our setting is considerably different, their optimal contract and our Pareto optimal equilibria exhibit similar dynamics: our account balance plays the same role in our model that real sunk investment plays in theirs.

Lastly, there is a connection between the present work and the literature on linked decisions. Jackson and Sonnenschein (2007) show that the ability to connect a large number

⁷Also see Abdulkadiroglu and Bagwell (2013), Nayyar (2009), Espino, Kozlowski, and Sanchez (2016), and Thomas and Worrall (1990) on the same theme. Our model features one-sided incomplete information like Thomas and Worrall (1990), but two-sided lack of commitment like the others.

⁸Also see Clementi and Hopenhayn (2006), DeMarzo and Fishman (2007), Edmans et al. (2012), and Malenko (2016).

of physically independent decisions helps align incentives.⁹ Frankel (2016b) shows that a principal with commitment power optimally employs a discounted quota to discipline an agent with state-independent preferences, and he shows that such quotas are strictly incentive compatible for the agent under partial alignment. In our model—without such commitment power—a form of dynamic budgeting is optimal but is tempered by the principal’s need for credibility.

Structure of the paper The remainder of the paper is structured as follows. Section 2 presents the model and introduces convenient language for discussing players’ incentives in our model. In Section 3, we discuss aligned equilibria, i.e. those in which no bad projects are adopted. Section 4 studies the dynamics of the relationship, showing that it becomes more conservative as time progresses. The heart of the paper is Section 5, wherein we fully characterize the equilibrium payoff set and the path of play of efficient equilibria. In Section 6, we present the Dynamic Capital Budget, a straightforward protocol that implements Pareto efficient equilibria. In Section 7, we discuss some possible extensions of our model. Final remarks follow in Section 8.

2 The Model

We consider an infinite-horizon game played between a principal ($\mathcal{P}$ principal) and an agent ($\mathcal{A}$ agent). Time is discrete and indexed by $k = 0, 1, 2, \dots$. The principal and the agent play the same stage game in every period.

Each period, the principal chooses whether or not to delegate a project adoption choice to the agent. That is, $\mathcal{P}$ publicly decides whether to *freeze* project adoption or to *delegate* it. If $\mathcal{P}$ freezes, then no project is adopted and both players accrue no payoffs that period. If $\mathcal{P}$ delegates, then $\mathcal{A}$ privately observes which type of project is available and decides whether or not to initiate the available project. The current period’s project is good (i.e. of type $\bar{\theta}$) with probability $h \in (0, 1)$ and bad (i.e. of type $\underline{\theta}$) with complementary probability. If the agent initiates a project of type θ , payoffs $(\theta - c, \theta)$ accrue to the players. Note that the cost is independent of the project’s type, and the difference between the agent’s payoffs and the principal’s payoffs does not depend on the agent’s private information. The principal observes whether or not the agent initiated a project, but she never sees the project’s type.

⁹Also see Casella (2005), Frankel (2016a), and Rubinstein (1979).

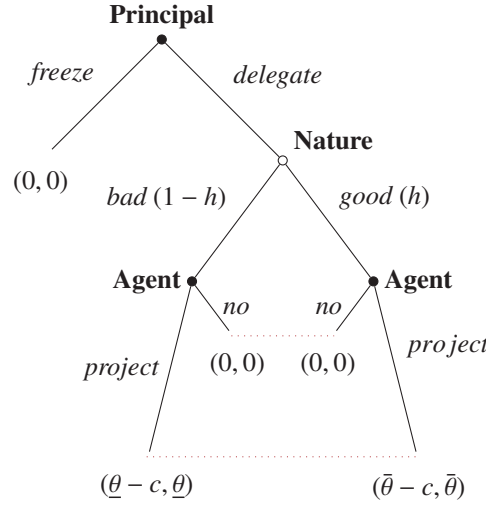


Figure 1: Stage game played by the principal and agent at each period. At the end of the period, the principal observes the agent's choice but not project quality.

To capture the preference misalignment between agent and principal in the stage game, we restrict our attention to the case that

$$0 < \underline{\theta} < c < \bar{\theta}.$$

First, given $\underline{\theta} - c < 0 < \bar{\theta} - c$, the principal prefers good projects to nothing, but prefers inactivity to bad projects. Second, given $0 < \underline{\theta} < \bar{\theta}$, the agent prefers any project to no project, but also prefers good ones to bad ones.

Notation. Let $\theta_E := (1 - h)\underline{\theta} + h\bar{\theta}$, the *ex-ante expected project value*.

For the preference misalignment between agent and principal to persist across periods, assume that the conflict of interest prevails ex-ante: $\theta_E < c$. That is, good projects are sufficiently rare that the principal would rather adopt no project than adopt every project. If the players interacted only once, the agent would not be selective. Indeed, the stage game has a unique sequential equilibrium: the principal freezes, and the agent takes a project if allowed.

While the players rank projects in the same way, the key tension in our model is a disagreement over which projects are worth taking. We interpret this payoff structure as the principal innately caring about the agent's (unobservable) payoff, in addition to the cost that she alone bears. The agent cares only about the benefit generated by a project, while

the principal cares about said benefit net of cost; we find revenue and profit to be useful interpretations of the players' payoffs.

The players share a common discount factor $\delta \in (0, 1)$, and they maximize **expected discounted profit** and **expected discounted revenue**, respectively. So, if the available project in each period $k \in \mathbb{Z}_+$ is θ_k , and projects are adopted in periods $\mathcal{K} \subseteq \mathbb{Z}_+$, then the principal and agent enjoy profit and revenue,

$$\pi = (1 - \delta) \sum_{k \in \mathcal{K}} \delta^k (\theta_k - c) \text{ and } v = (1 - \delta) \sum_{k \in \mathcal{K}} \delta^k \theta_k, \text{ respectively.}$$

Before proceeding to characterize the equilibrium payoff set, we document one more piece of notation, and we make an interpretable assumption which will guarantee existence of interesting equilibria.

Notation. Let $\omega := h(\bar{\theta} - \underline{\theta}) = \theta_E - \underline{\theta}$, the *marginal value of search*.

The constant ω captures the marginal option value of seeking another project rather than taking an existing bad project.

Assumption 1.

$$\delta\omega \geq (1 - \delta)\underline{\theta} \text{ or, equivalently, } \omega \geq (1 - \delta)\theta_E.$$

Henceforth, we take Assumption 1 as given. This assumption, which can equivalently be expressed as a lower bound on the discount factor δ , highlights the power of intertemporal tradeoffs to incentivize good behavior from the agent. If the agent is sufficiently patient, the marginal value of searching for a good project outweighs the myopic benefit of an immediate bad project.

While the misalignment of preferences (in particular, $c > \theta_E$) tells us that the only *stationary* equilibrium of our repeated game is an unproductive one—i.e. never has any (good or bad) projects adopted—Assumption 1 guarantees existence of productive equilibria. Indeed, it is necessary and sufficient for the following to be an equilibrium: before the first project is adopted, the principal delegates and the agent adopts only good projects; after the first project is adopted, they play the stage game equilibrium. Assumption 1 exactly delivers the agent's incentive constraint to willingly resist a bad project.

Lastly, we define another lower bound on the discount factor.

Notation. Let $\bar{\delta} := \frac{\bar{\theta}(c - \theta_E)}{\bar{\theta}(c - \underline{\theta}) - c(\theta_E - \underline{\theta})}$, the *players' patience threshold*.

When the discount factor is at least $\bar{\delta}$, we will show that there exist equilibria with initial project adoption.¹⁰ This case yields rich dynamics for the relationship.

Expressing Payoffs Every strategy profile entails an expected discounted number of adopted good projects $g = (1 - \delta)\mathbb{E} \sum_{k \in \mathcal{K}} \delta^k \mathbf{1}_{\{\theta_k = \bar{\theta}\}}$ and an expected discounted number of adopted bad projects $b = (1 - \delta)\mathbb{E} \sum_{k \in \mathcal{K}} \delta^k \mathbf{1}_{\{\theta_k = \underline{\theta}\}}$, where $\mathcal{K} \subseteq \mathbb{Z}_+$ is the realized set of periods in which the principal delegates and the agent adopts a project. Given those, one can compute the agent value/revenue as

$$v = \bar{\theta}g + \underline{\theta}b$$

and the principal value/profit as

$$\pi = (\bar{\theta} - c)g - (c - \underline{\theta})b.$$

For ease of bookkeeping, it is convenient to track equilibrium-supported revenue v and bad projects b , both in expected discounted terms. The vector (v, b) encodes both agent value v and principal profit,

$$\begin{aligned} \pi(v, b) &:= (\bar{\theta} - c)g - (c - \underline{\theta})b \\ &= (\bar{\theta} - c) \frac{v - \underline{\theta}b}{\bar{\theta}} - (c - \underline{\theta})b \\ &= \left(1 - \frac{c}{\bar{\theta}}\right)v - c \left(1 - \frac{\underline{\theta}}{\bar{\theta}}\right)b. \end{aligned}$$

In (v, b) space, the principal's indifference curves have constant slope $m_0 := \frac{\bar{\theta} - c}{c(\bar{\theta} - \underline{\theta})}$.

2.1 Equilibrium Values

To understand the payoffs that can be attained in equilibrium, we proceed by defining an equilibrium concept and using the specific structure of our game to characterize the equilibrium payoff set.

Throughout the paper, *equilibrium* will be taken to mean perfect public equilibrium (PPE),¹¹ in which the players' stage game strategies respond only to the history of public

¹⁰This is not a patient limit result. Holding $(\bar{\theta}, \underline{\theta}, h)$ fixed, $\bar{\delta}$ runs from 0 to 1 as c ranges from θ_E to $\bar{\theta}$.

¹¹This is, in a sense, without loss. As we will show that all equilibrium payoffs are attainable in pure

outcomes and public randomization devices.

Definition 1. *Each period, one of three public outcomes occurs: the principal freezes; the principal delegates and the agent initiates no project; or the principal delegates and the agent initiates a project. A period- k **public history**, h_k , is a sequence of k public outcomes (along with realizations of public randomization devices for each period).*

*A **public principal strategy** specifies, for each public history, a probability of delegation. A **public agent strategy** specifies, for each public history, a project type-contingent probability of project adoption.*

*A **perfect public equilibrium** (PPE) is a sequential equilibrium in which both players use public strategies.*

Toward a Characterization The main objective of this paper is to characterize the set of equilibrium-supported payoffs,

$$\mathcal{E}^* := \{(v, b) : \exists \text{ equilibrium with revenue } v \text{ and bad projects } b\} \subseteq \mathbb{R}_+^2.$$

As neither the principal nor the agent can commit to a strategy (i.e. sequential rationality is required for both), we must rely on the specific structure of our game to characterize the limits of the relationship and, in particular, which payoffs can be supported in equilibrium.

Throughout the paper, we make extensive use of two simple observations about our model. First, notice that $(0, 0) \in \mathcal{E}^*$, since the profile σ^{static} , in which the principal always freezes and the agent takes every permitted project, is an equilibrium. Said differently, there is always an unproductive equilibrium—i.e. one with no projects. This equilibrium provides min-max payoffs to both players, so that they can only benefit from their recurring interaction. Second, as the following lemma clarifies, off-path strategy specification is unnecessary in our model. If appropriate on-path incentive constraints are satisfied, we can alter behavior off-path to yield an equilibrium. With the lemma in hand, we rarely specify off-path behavior in a given strategy profile.

Lemma 1. *Fix a strategy profile σ , and suppose that:*

1. *The agent has no profitable deviation from any on-path history.*
2. *At all on-path histories, the principal has nonnegative continuation profit.*

strategies, standard results (see Mailath and Samuelson (2006)) tell us that every sequential equilibrium has a payoff-equivalent equilibrium in public strategies.

Then, there is an equilibrium $\tilde{\sigma}$ that generates the same on-path behavior (and, therefore, the same value profile).

Proof. Let σ^{static} be the stage Nash profile—i.e. the principal always freezes, and the agent takes a project immediately whenever permitted. Define $\tilde{\sigma}$ as follows:

- On-path (i.e. if $\mathcal{P}$ has never deviated from σ), play according to σ .
- Off-path (i.e. if $\mathcal{P}$ has ever deviated from σ), play according to σ^{static} .

The new profile is incentive-compatible for the agent: off-path because σ^{static} is, on-path because σ is. It is also incentive-compatible for the principal: off-path because σ^{static} is, on-path because σ is and has nonnegative continuation profits while σ^{static} yields zero profit. $\square$

Dynamic Incentives We proceed by characterizing the incentives of both players in our dynamic setting. To understand the players' incentives at any given moment, we must first understand how their future payoffs respond to their current choices. To describe the law of motion of revenue v [or, respectively, bad projects b], we keep track of continuation values, conditional on public outcomes:

- v_F , next period's agent continuation value if the principal freezes today;
- v_N , the agent continuation value if the principal delegates today and the agent abstains from project adoption;
- v_P , the agent continuation value if a project is undertaken today; and
- b_F, b_N, b_P defined analogously for continuation (expected discounted) bad projects.

Continuation values cannot depend on the quality of an adopted project, nor on the quality of a forgone project, which are not publicly observable. Finally, we describe the players' present actions (after any public randomization occurs) as follows:

- The principal makes a delegation choice $p \in [0, 1]$, the probability with which she delegates to the agent in the current period.
- The agent chooses $\bar{a} \in [0, 1]$ and $a \in [0, 1]$, the probabilities which he currently initiates available good and bad projects, respectively, conditional on being allowed to.

Appealing to self-generation arguments, as in Abreu, Pearce, and Stacchetti (1990), and to Lemma 1, equilibrium is characterized by the following three conditions:

1. Promise keeping:

$$(v, b) = ph\bar{a} \left[(1 - \delta)(\bar{\theta}, 0) + \delta(v_P, b_P) \right] + p(1 - h)a \left[(1 - \delta)(\underline{\theta}, 1) + \delta(v_P, b_P) \right] \\ + p[h(1 - \bar{a}) + (1 - h)(1 - a)]\delta(v_N, b_N) + (1 - p)\delta(v_F, b_F).$$

We decompose continuation outcomes (v, b) after any period into what happens in each of four events: (i) the agent finds and invests in a good project; (ii) the agent finds and invests in a bad project; (iii) no project is adopted by the agent's choice; and (iv) no project is adopted on the principal's authority. Each weighted by their equilibrium probabilities and subject to the restriction that continuation values respond only to public information.

2. Agent incentive compatibility:

$$a \in \operatorname{argmax}_{\hat{a} \in [0,1]} \hat{a} \left[(1 - \delta)\underline{\theta} + \delta v_P \right] + (1 - \hat{a})\delta v_N, \\ \bar{a} \in \operatorname{argmax}_{\hat{a} \in [0,1]} \hat{a} \left[(1 - \delta)\bar{\theta} + \delta v_P \right] + (1 - \hat{a})\delta v_N.$$

The agent chooses whether or not to invest in a project by balancing the immediate gain with the changes in her continuation payoff. In particular, if the agent is willing to resist taking a project immediately ($a < 1$), it must be that the punishment in continuation payoff for taking a project, $\delta(v_N - v_P)$, is severe enough to deter the $(1 - \delta)\underline{\theta}$ myopic gain.

3. Principal participation:¹²

$$\pi(v, b) \geq 0.$$

The principal could, at any moment, unilaterally move to a permanent freeze and secure herself a profit of zero. Therefore, at any history, she must be securing at least that much in equilibrium.

¹²This is, of course, a relaxation of the principal's incentive constraint. Even in light of Lemma 1, the principal must be indifferent between freezing and delegating for $p \in (0, 1)$ to be incentive-compatible. In the appendix, we show that it is without loss for the principal to never privately mix. With this in mind, the present relaxed incentive constraint is all that matters.

3 Aligned Equilibrium

Our game has no productive stationary equilibrium. If the principal allows history-independent project adoption, the agent cannot be stopped from taking limitless bad projects. In the present section, we ask to what extent this core tension can be resolved by allowing non-stationary equilibria.

Definition. An *aligned equilibrium* is an equilibrium in which no bad projects are ever adopted on-path.

As an example, consider the following straightforward strategy profile σ_τ , indexed by a number of freeze-periods, $\tau \in \mathbb{Z}_+$. The principal initially delegates, and the agent adopts the first project if and only if it is of type $\bar{\theta}$. If the agent abstains today, σ_τ begins again tomorrow. If the agent takes a project today, then the principal freezes for τ periods. After τ periods of no delegation, σ_τ begins again.

When the principal delegates to the agent, the agent always serves the principal's interests. The freeze of length τ is used to incentivize the agent to exert restraint. If and only if the freeze duration is high enough to make an unlucky (i.e. facing a $\underline{\theta}$ project) agent exercise restraint—which Assumption 1 guarantees will happen for high enough τ —this profile will be an equilibrium. The optimal equilibrium in this class leaves the incentive constraint $\delta(v - \delta^\tau v) \geq (1 - \delta)\underline{\theta}$ binding,¹³ as lowering τ yields a Pareto improvement. For this minimal choice of punishment,

$$v = (1 - \delta)h\bar{\theta} + \delta h(v_P - v) + \delta v = \delta v + h(1 - \delta)(\bar{\theta} - \underline{\theta}) \implies v = \omega.$$

The agent's continuation value whenever the principal delegates is exactly equal to his marginal value of search.

This simple class illuminates the forces at play in our model. The principal wants many good projects to be initiated, but she cannot give the agent free rein. If she wants to stop him from investing in bad projects, she must threaten him with mutual surplus destruction. Subject to wielding a large enough stick to encourage good behavior, she efficiently wastes as little opportunity as possible.

This sensible equilibrium is in fact optimal among all aligned equilibria.

¹³In the case that such τ is not an integer, we can use public randomization between two integer values of τ to make $\delta[v - \mathbb{E}(\delta^\tau v)] = (1 - \delta)\underline{\theta}$. Under Assumption 1, the intermediate value theorem tells us such a stochastic τ exists.

Proposition 1 (Aligned Optimum). *The highest agent value in any aligned equilibrium is ω , the marginal value of search.*¹⁴

The intuition behind the proposition, formally proven in the appendix, is straightforward. If the agent’s current expected value exceeds his marginal value of search, then asking an agent to wait (i.e. to exercise restraint when facing a bad project) requires growth in his continuation value. Eventually, after enough bad luck, this promised value can only be fulfilled by allowing some bad projects.

Corollary 1. *The aligned optimal profit is $\omega \left(1 - \frac{c}{\bar{\theta}}\right)$, which is independent of players’ patience and strictly below the principal’s first-best profit $h(\bar{\theta} - c)$.*

To keep the agent honest, the principal must destroy some future surplus when the agent overspends. As the players become more patient, the punishment must be made more draconian to provide the same incentives. When only good projects are adopted, this added severity punishes the principal too, so that no welfare gains are achieved.

We know there is an equilibrium providing the principal with payoffs near her first-best—and thus Pareto dominating all aligned equilibria, by Corollary 1—if the players are sufficiently patient.¹⁵ The corollary therefore tells us that aligned equilibria are inadequate: bad projects are a necessary ingredient of a well-designed relationship.

4 Relationship Dynamics

The previous section showed us that a relationship of ongoing delegation should, at least for some parameter values, admit some bad projects. This complicates analysis in two ways. First, there is an extra degree of freedom in providing agent value: not just when to have the agent adopt projects, but also which types to have him adopt. Second, the principal’s incentives are now relevant. If, after some history, the agent plans to adopt too many bad projects, the principal will strictly prefer to unilaterally freeze.

Our ultimate goal, which we will achieve in Section 5, is to characterize Pareto efficient equilibria and the value they produce. In the present section, we derive some key qualitative properties of the path of play of such equilibria.

¹⁴As for all results in the main text, this result holds under Assumption 1. Absent Assumption 1, there can be no aligned equilibrium yielding $v > 0$.

¹⁵While this follows from our later results (e.g. Theorem 1), it could alternatively be deduced from a folk theorem. Indeed, by first removing two dominated stage game $\mathcal{A}$ strategies (i.e. those in which he does not adopt good projects), the results of Fudenberg, Levine, and Maskin (1994) apply.

The next theorem describes the dynamics of an efficient relationship between the two players. In the early stages of the relationship, every good project is adopted, but some bad ones are as well: the per-period revenue is high, but the average project taken in this stage has a return on investment strictly below $\frac{\bar{\theta}}{c}$. As the relationship progresses and stochastic project types unfold, the regime eventually changes. In the later stages of the relationship, some good projects are missed, but no bad projects are adopted: the per-period revenue is now lower, but the average project has a higher return on investment. Said differently, the players transition from a high-revenue, low-yield relationship to a low-revenue, high-yield relationship.

Theorem 1 (Relationship Dynamics). *In any Pareto efficient equilibrium, there is (with probability 1) a finite time K such that:*

- *In period $k < K$: the agent's continuation value strictly exceeds ω , the principal delegates, and the agent adopts the current project if it is good.*
- *In period $k \geq K$: The agent's continuation value is weakly below ω , and the agent does not adopt the current project if it is bad.*

Moreover, if $\delta \geq \bar{\delta}$, then $K > 0$.

Proof. We prove the theorem through a series of smaller claims. First note that $\bar{v} := \max\{v : (v, b) \in \mathcal{E}^*\} < \theta_E$. Indeed, the unique b such that (θ_E, b) is feasible is $b = 1 - h$, yielding negative profit to the principal. Now, let $\epsilon := \frac{1-\delta}{\delta} \left(\frac{\theta_E - \bar{v}}{2} \right) > 0$ and $q := \frac{\epsilon}{\bar{v} + \epsilon} \in (0, 1)$.

Claim 1: If the agent's continuation value is v before a public randomization, then the value $\hat{v}$ after public randomization is weakly below $v + \epsilon$ with probability at least q .

Proof: Letting $\hat{q} := \mathbb{P}\{\hat{v} \leq v + \epsilon\}$ yields $v = (1 - \hat{q})\mathbb{E}\{\hat{v} | \hat{v} > v + \epsilon\} + \hat{q}\mathbb{E}\{\hat{v} | \hat{v} \leq v + \epsilon\} \geq (1 - \hat{q})(v + \epsilon) + \hat{q}0 = (v + \epsilon) - \hat{q}(v + \epsilon)$, so that $\hat{q} \geq \frac{\epsilon}{v + \epsilon} \geq q$.

Claim 2: If (i) $(v, b), (v^1, b^1)$ are equilibrium payoffs with $(v, b) \in \text{co}\{(v^1, b^1), (0, 0)\}$,¹⁶ (ii) (v, b) is some on-path continuation value pair for some Pareto efficient equilibrium; and (iii) $v \geq \omega$; then $(v, b) = (v^1, b^1)$.

Proof: If not, then $(v, b) = (\lambda v^1, \lambda b^1)$ for some $\lambda \in (0, 1)$. There is a unique $(v_N, b_N) \in \text{co}\{(v^1, b^1), (\omega, 0)\}$ with $v_N = v$. By construction, $b_N < b$, so that the equilibrium— $\mathcal{E}^*$ being convex—value pair (v_N, b_N) provides the exact same agent value as (v, b) but at a strictly higher principal value. Replacing the continuation play with some continuation equilibrium

¹⁶We use “co” to refer to the convex hull of a set.

providing (v_N, b_N) maintains equilibrium, contradicting Pareto optimality of the original equilibrium.

Claim 3: If the agent's continuation value is $v \geq \omega$ after public randomization on the path of some Pareto efficient equilibrium, then the next period's continuation value (before public randomization) following a good project today is weakly exceeded by $v - 2\epsilon$.

Proof: Suppose the agent's continuation value is $v \geq \omega$ after public randomization; let b be the expected discounted number of bad projects. Let stage play $p, \bar{a}, a$, and continuation payoffs $v_F, v_N, v_P, b_F, b_N, b_P$ be as described in Section 2.

By convexity, $(v, b) = (1-p)(1-\delta)(0, 0) + [1-(1-p)(1-\delta)](v^1, b^1)$ for some $(v^1, b^1) \in \mathcal{E}$. Therefore, $p = 1$ by Claim 2; the principal must be delegating.

Now, if $\bar{a} < 1$ then agent IC would tell us $a = 0$, so that

$$\begin{aligned} (v, b) &= h \left\{ \bar{a} \left[(1-\delta)(\bar{\theta}, 0) + \delta(v_P, b_P) \right] + (1-\bar{a})\delta(v_N, b_N) \right\} + (1-h)\delta(v_N, b_N) \\ &= (1-\bar{a})\delta(v_N, b_N) + \bar{a} \left\{ h \left[(1-\delta)(\bar{\theta}, 0) + \delta(v_P, b_P) \right] + (1-h)\delta(v_N, b_N) \right\} \\ &\in (1-\bar{a})\delta(v_N, b_N) + \bar{a}\mathcal{E}^*. \end{aligned}$$

To see the last containment, consider two cases. If $\bar{a} = 0$, there is nothing to show. If $\bar{a} \in (0, 1)$, then Pareto optimality (relative to an alternative mixture) requires that

$$\pi \left(h \left[(1-\delta)(\bar{\theta}, 0) + \delta(v_P, b_P) \right] + (1-h)\delta(v_N, b_N) \right) = \pi(\delta(v_N, b_N)),$$

in which case $h \left[(1-\delta)(\bar{\theta}, 0) + \delta(v_P, b_P) \right] + (1-h)\delta(v_N, b_N) \in \mathcal{E}^*$.

Again by Claim 2, it cannot be that $(v, b) \in (1-\bar{a})\delta(v_N, b_N) + \bar{a}\mathcal{E}^*$. So $\bar{a} = 1$; the agent takes every good project. Hence,

$$\begin{aligned} v &\geq (1-\delta)\theta_E + \delta v_P, \text{ by agent IC.} \\ \implies \delta v - \delta v_P &\geq (1-\delta)\theta_E - (1-\delta)v \\ \implies v - v_P &\geq 2\epsilon, \text{ proving the claim.} \end{aligned}$$

Claim 4: In any Pareto efficient equilibrium, there is (with probability 1) some period K such that the agent's continuation value is weakly exceeded by ω .

Proof: Combining Claims 1 and 3, there is a probability of at least qh of the continuation value decreasing by at least ϵ in a given period. Therefore, for any $k \in \mathbb{N}$ and continuation value in $[\omega, \omega + k\epsilon]$, there is a probability of at least $(qh)^k$ of the continuation value de-

creasing below ω in the next k periods. As $[0, \bar{v}]$ is bounded, the continuation value almost surely falls below ω at some time.

Claim 5: If the agent's continuation value is $v \leq \omega$ at some on-path history of a Pareto efficient equilibrium, his value is less than or equal to ω forever.

Proof: Let b be the expected discounted number of bad projects from the current history. Appealing to Proposition 1, it must be that $b = 0$. Otherwise, we could replace the continuation play with a continuation equilibrium generating $(v, 0)$ and preserve equilibrium. Therefore, by promise-keeping, it must be that every on-path future history (v^1, b^1) has $b^1 = 0$ as well. Again by Proposition 1, $v^1 \leq \omega$.

Together, Claims 4 and 5 show that such a finite time K exists. Finally, in the case that $\delta \geq \bar{\delta}$, Lemma 13 in the appendix tells us that the agent's ex-ante expected payoff is strictly higher than ω in some equilibrium; while Lemma 12 tells us that B is differentiable at ω , and therefore $\max_{v \in [0, \omega]} \pi(v, B(v)) = \pi(\omega, B(\omega)) < \pi(\omega + \epsilon, B(\omega + \epsilon))$ for sufficiently small $\epsilon > 0$. Therefore, any Pareto optimal equilibrium has an initial agent value strictly above ω , i.e. $K > 0$. $\square$

From the agent's perspective, the early stages of the relationship are unequivocally better than the later stages; his freedom declines over time. From the principal's perspective, the profit implications are ambiguous; she wants to produce a high agent revenue as in the early stages, but she also wants a high yield as in the later stages.

5 Characterizing Equilibrium

We now know that the delegation relationship transitions toward conservatism in some (possibly distant) future, but this is an incomplete description. How much productive delegation happens in the early, more flexible phase? When will bad projects be financed? While Theorem 1 tells us a lot about the trajectory of the relationship, it is silent on these more detailed questions about the nature of repeated delegation.

In this section, we explicitly describe the equilibrium value set of our repeated game—our main theorem. From this, we then derive some key implications for behavior in Pareto optimal equilibria.

5.1 The Equilibrium Value Set

First, we offer a complete characterization of the equilibrium value set. The heart of the proof is the characterization of the equilibrium payoff frontier B , formally carried out in the appendix. The overall structure of the argument proceeds as follows. First, allowing for public randomization guarantees convexity of the equilibrium payoff frontier. As a consequence, whenever the agent is to adopt only good projects in some period, the principal optimally inflicts the minimum punishment that provides appropriate agent incentives. Next, because the principal has no private action or information, we show that the frontier is self-generating and private mixing unnecessary. Finally, we show that initial freeze is inefficient for values above ω , and that bad project adoption is wasteful, except when used to provide very high agent values.

Recall that $m_0 = \frac{\bar{\theta}-c}{c(\bar{\theta}-\underline{\theta})}$, the slope of the zero-profit line.

Theorem 2 (Equilibrium Value Set). *There is a value $\bar{v} \geq \omega$ and a continuous function $B : [0, \bar{v}] \rightarrow \mathbb{R}_+$ such that $\mathcal{E}^* = \{(v, b) \in [0, \bar{v}] \times \mathbb{R} : B(v) \leq b \leq m_0 v\}$, and:*

- $B|_{[0, \omega]} = 0$.
- On $[\omega, \delta\bar{v} + (1-\delta)\omega]$, B is strictly convex with

$$B(v) = \delta \left[(1-h)B\left(\frac{v - (1-\delta)\omega}{\delta}\right) + hB\left(\frac{v - (1-\delta)\omega - (1-\delta)\underline{\theta}}{\delta}\right) \right].$$

- On $[\delta\bar{v} + (1-\delta)\omega, \bar{v}]$, B is affine.

Moreover, if $\delta \geq \bar{\delta}$, then: $\bar{v} > \omega$, $B(\bar{v}) = (1-\delta)(1-h) + \delta B\left(\frac{v - (1-\delta)\omega - (1-\delta)\underline{\theta}}{\delta}\right)$, and $\pi(\bar{v}, B(\bar{v})) = 0$.

It may at first appear surprising that one simple scalar equation can describe the full equilibrium set, even implicitly. The driving observation is that, with no ex-post monitoring of project types, the agent's binding incentive constraint completely pins down the law of motion for the agent's continuation value.

With $\mathcal{E}^*$ in hand, its Pareto frontier and the principal-optimal equilibrium value are immediate.

Corollary 2. *Let $\bar{v}, B$ be as in the statement of Theorem 2, and suppose $\delta \geq \bar{\delta}$. There is then a value $v^* \in (\omega, (1-\delta)\bar{v} + \delta\omega]$ such that:*

- The uniquely profit-maximizing equilibrium value profile is $(v^*, B(v^*))$.

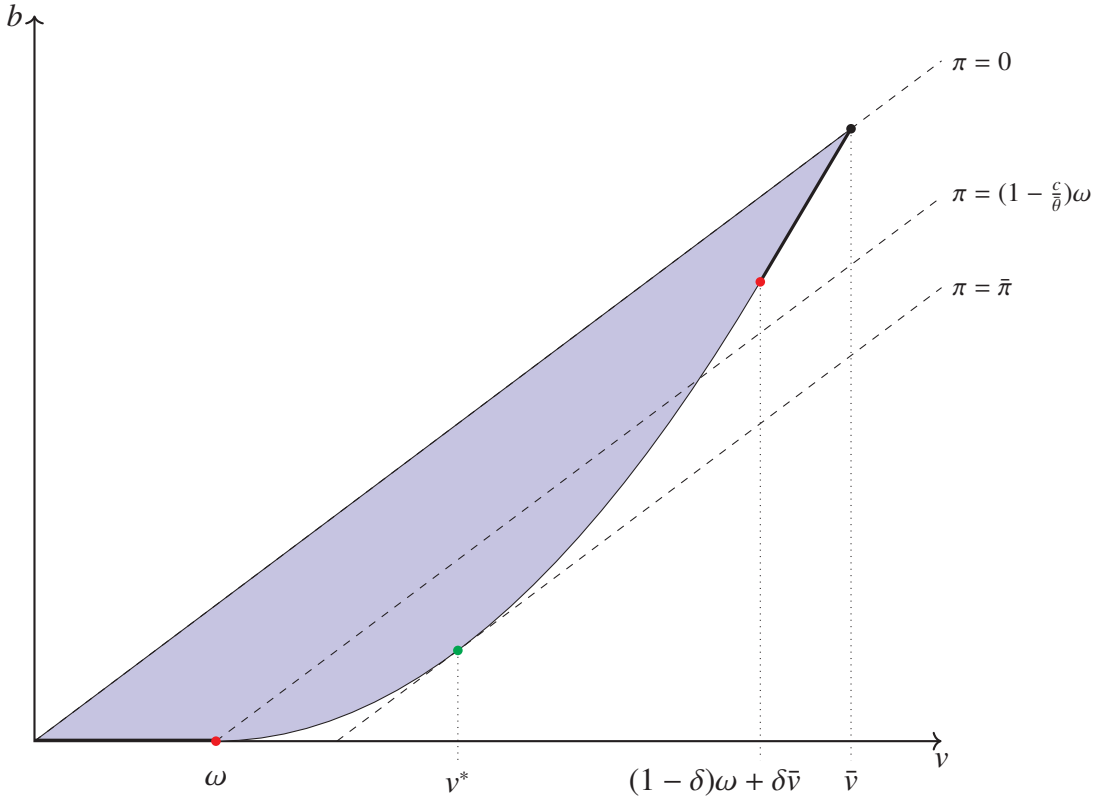


Figure 2: The solid line traces out the frontier B . The equilibrium value set is the convex region between B and the dashed zero-profit line. The dashed lines trace different isoprofits. The green dot highlights the uniquely principal-optimal vector, $\bar{\pi}$.

- The Pareto frontier of $\mathcal{E}^*$ is $\{(v, B(v)) : v^* \leq v \leq \bar{v}\}$.

Proof. Let $\pi^* : [0, \bar{v}] \rightarrow \mathbb{R}$ be given by $\pi^*(v) = \pi(v, B(v))$. Then $\operatorname{argmax}_v \pi^*(v)$ is a subset of $[\omega, \bar{v}]$ because $\pi^*|_{[0, \omega]}$ is strictly increasing; a subset of $[\omega, (1 - \delta)\bar{v} + \delta\omega]$ because $\pi^*|_{[(1 - \delta)\bar{v} + \delta\omega, \bar{v}]}$ is strictly decreasing;¹⁷ nonempty because π^* is continuous; and a singleton (say v^*) because $\pi^*|_{[\omega, (1 - \delta)\bar{v} + \delta\omega]}$ is strictly concave. The frontier characterization follows because π^* is increasing on $[0, v^*]$ and decreasing on $[v^*, \bar{v}]$.

It remains to verify that $v^* \neq \omega$. By Lemma 12 in the appendix, B is differentiable at ω with $B'(\omega) = 0$. Therefore, π^* is differentiable at ω with derivative $1 > 0$ and so cannot be maximized at ω . $\square$

¹⁷By Theorem 2, the function B is affine on $[(1 - \delta)\bar{v} + \delta\omega, \bar{v}]$, and so π^* is as well. As π^* is nonnegative, concave, and (again by Theorem 2) zero at its endpoints if $\delta \geq \bar{\delta}$, there are only two possibilities: π^* is the constant zero function, or π^* is strictly decreasing on $[(1 - \delta)\bar{v} + \delta\omega, \bar{v}]$. As $\pi^*(\omega) = (1 - \frac{\epsilon}{\theta})\omega > 0$, the latter holds.

5.2 Features of Equilibrium Behavior

We are now, with the equilibrium payoff set in hand, equipped to say more about the players' behavior under Pareto efficient equilibria. As Theorem 2 and its Corollary 2 show, every Pareto efficient equilibrium value (with agent value $v_0 \in [v^*, \bar{v}]$) can be implemented in an equilibrium with the features described in the following definition.

Definition 2. Fix an equilibrium σ , and define the associated stochastic process $\{v_k, \hat{v}_k\}_{k=0}^\infty$, where v_k [resp. $\hat{v}_k$] is the agent's time k continuation before [after] that period's public randomization.

Define the following features of an equilibrium:

- Say σ has **no superfluous randomization** if:

$$\begin{aligned} v_k \in [\omega, \delta\bar{v} + (1 - \delta)\omega] \cup \{\bar{v}\} &\implies \hat{v}_k = v_k \text{ with probability 1;} \\ v_k \in (\delta\bar{v} + (1 - \delta)\omega, \bar{v}) &\implies \hat{v}_k \in [\delta\bar{v} + (1 - \delta)\omega, \bar{v}] \text{ with probability 1.} \end{aligned}$$

- Say σ has **good behavior with minimal punishment** if:

$$\hat{v}_k \in [\omega, \delta\bar{v} + (1 - \delta)\omega] \implies \begin{cases} \mathcal{P} \text{ delegates;} \\ \mathcal{A} \text{ adopts the project if and only if it is good;} \\ v_{k+1} = \frac{\hat{v}_k - (1 - \delta)\omega}{\delta} \text{ if the agent does not adopt project } k; \\ v_{k+1} = \frac{\hat{v}_k - (1 - \delta)\omega - (1 - \delta)\theta}{\delta} \text{ if the agent adopts project } k. \end{cases}$$

- Say σ has **gift-giving as a last resort** if:

$$\hat{v}_k > (1 - \delta)\omega + \delta\bar{v} \implies \begin{cases} \mathcal{P} \text{ delegates;} \\ \mathcal{A} \text{ adopts project } k; \\ v_{k+1} = \frac{\hat{v}_k - (1 - \delta)\omega - (1 - \delta)\theta}{\delta}. \end{cases}$$

- Say σ has **absorbing punishment** if:

$$v_k \leq \omega \implies v_\ell, \hat{v}_\ell \leq \omega \quad \forall \ell \geq k \text{ with probability 1.}$$

To gain some intuition as to why the above should be features of efficient play, consider how the principal might like to provide different levels of revenue. The case of revenue

$v \leq \omega$ is simple: Proposition 1 tells us that the principal can provide said revenue efficiently, via aligned equilibrium. The case of $v = \bar{v}$ is straightforward too: to keep her current promise to the agent, the principal is forced to allow a project today, no questions asked. The described form of B between the two obtains from having the agent defend the principal's interests *whenever possible*. Reminiscent of Ray (2002), costly incentive provision is backloaded as much as possible. The principal's limited commitment constrains this backloading, capping how much revenue can be provided while asking the agent to defend the principal's interests.

Uniqueness As the next result shows, this limited backloading is a necessary feature of any Pareto optimal equilibrium. That is, the path of play implicit in Theorem 2—until the absorbing, conservative phase of the relationship is reached—is essentially unique.¹⁸

Corollary 3. *Every Pareto efficient equilibrium entails (i) no superfluous randomization, (ii) good behavior with minimal punishment, (iii) gift-giving as a last resort, and (iv) absorbing punishment.*

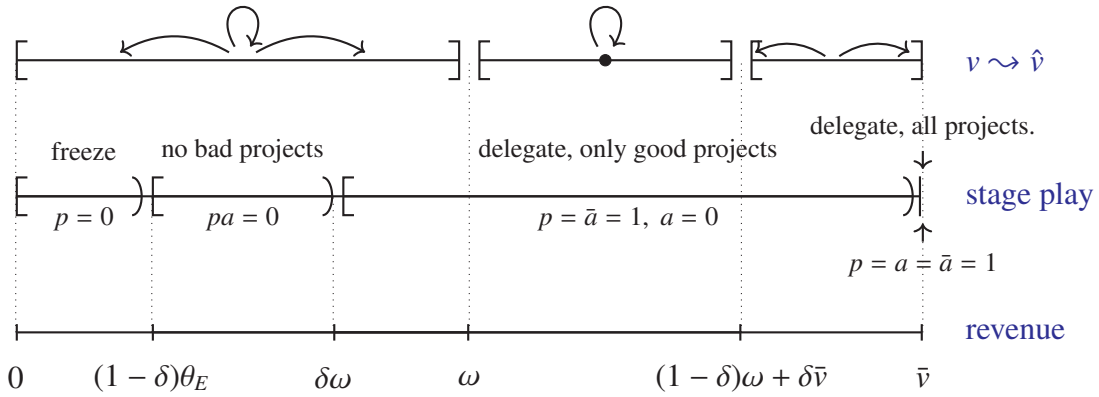


Figure 3: Necessary features of Pareto optimal equilibrium.

Proof. First recall that any Pareto efficient equilibrium will never have continuation value strictly above the graph of B at any on-path history. Indeed, replacing the continuation

¹⁸“Essentially” because immediate public randomization can be used over $(\delta\bar{v} + (1-\delta)\omega, \bar{v})$ but may not be required at values near $\bar{v}$, if δ is small enough. This region vanishes in the continuous time limit we consider in Section 6.

value with one yielding the same agent value and strictly fewer bad projects will preserve equilibrium and yield a Pareto improvement.

With this in mind, (iii) and (iv) follow directly from promise keeping, and (i) and half of (ii)—namely, minimum punishment—follow from strict convexity of $B|_{[\omega, \delta\bar{v} + (1-\delta)\omega]}$. We only need to show that the agent adopts only good projects whenever $\hat{v} \in (\omega, \delta\bar{v} + (1-\delta)\omega)$.

For values of at least ω , freeze is never optimal. Now, Theorem 2 tells us

$$\begin{aligned} D(v) &:= \delta \left[(1-h)B\left(\frac{v-(1-\delta)\omega}{\delta}\right) + hB\left(\frac{v-(1-\delta)\omega-(1-\delta)\theta}{\delta}\right) \right] \\ &\quad - \left[(1-\delta)(1-h) + \delta B\left(\frac{v-(1-\delta)\omega-(1-\delta)\theta}{\delta}\right) \right] \\ &= (1-h)\delta \left[B\left(\frac{v-(1-\delta)\omega}{\delta}\right) - B\left(\frac{v-(1-\delta)\omega-(1-\delta)\theta}{\delta}\right) \right] - (1-h)(1-\delta) \end{aligned}$$

is strictly increasing in $v \in [\omega, \delta\bar{v} + (1-\delta)\omega]$ (since B is convex, strictly so around $\frac{v-(1-\delta)\omega-(1-\delta)\theta}{\delta}$) and globally nonpositive there, since asking the agent to adopt only good projects is optimal. Consequently, $D(v) < 0$ for $v < \delta\bar{v} + (1-\delta)\omega$. Therefore, unrestrained project adoption will not take place on the path of a Pareto efficient equilibrium while the agent's continuation value is in $[\omega, \delta\bar{v} + (1-\delta)\omega]$. Next, the agent must adopt all good projects there (by Theorem 1). Finally, Lemma 5 in the appendix tells us that the agent never mixes over good project adoption in Pareto optimal equilibrium, so that no bad projects are adopted. Point (ii) follows. $\square$

Figure 3 highlights the features of Pareto efficient equilibrium behavior codified in Corollary 3. The bottom line depicts the current promised revenue to the agent. The middle line describes features of stage play as a function of this promised revenue. In particular, note that for revenues higher than $\delta\omega$, efficient equilibrium requires that the principal delegates with certainty. If the revenue is not at its highest value, the agent then exercises restraint. Finally, the top line illustrates the impact of public randomization. When the promised revenue lies in $[\omega, (1-\delta)\omega + \delta\bar{v}]$, efficient equilibrium prescribes that public randomization is not used. Given the benefits of backloading, any needless noise in the agent's continuation value is strictly costly.

Bad Projects At the end of Section 3, we saw that bad projects would, for some parameter values, be a necessary part of Pareto efficient equilibrium. Initially, one might have suspected that the choice of whether or not to employ bad projects to provide incentives the the agent amounts to evaluating a profit tradeoff by the principal. Theorem 2 tells us

that the principal faces no such tradeoff. Whenever a promise of future bad projects can be credible,¹⁹ it is a necessary component of an optimal contract.

Corollary 4. *If $\delta \geq \bar{\delta}$, then every Pareto efficient equilibrium entails some bad projects.*

This follows directly from Corollary 2 (in particular, that $v^* > \omega$). Intuition for this result comes from considering equilibria as described in Corollary 3, with agent value $v = \omega + \epsilon$ strictly higher than, but very close to, the marginal value of search. The principal, with objective proportional to $m_0 v - b$, benefits from the higher agent value $\omega + \epsilon > \omega$. Corollary 3 tells us that the cost this entails—a bounded discounted number of bad projects—will only accrue in an exponentially distant future. Thus the principal strictly prefers an equilibrium with $\omega + \epsilon$ to the aligned optimum.

The takeaway is straightforward: if bad projects can be credibly promised in equilibrium, then they unambiguously should be. The principal should offer them to the agent later to sustain better incentives today, prolonging the benefits of delegation.

6 Implementation: Dynamic Capital Budgeting

There are efficiency gains to be had from allowing bad projects, but one must carefully balance the principal’s credibility constraint for it to remain an equilibrium. The implementation we discuss in this section, the **Dynamic Capital Budgeting** (DCB) contract, is an attempt to achieve this balance.

6.1 A Continuous Time Limit

For the purposes of this section, we find it convenient to work with a (heuristic) continuous time limit of our game in which the players interact very frequently, but good projects remain scarce. Letting the time between decisions, together with the proportion of good projects, vanish enables us to present our budget rule cleanly, in the language of calculus.

Suppose the players discount time $t \in \mathbb{R}_+$ at a rate of $r > 0$ and good projects arrive with Poisson rate $\eta > 0$. Assume bad projects are abundant, in the sense that (regardless of history) a bad project is available to adopt whenever a good project is not. The players meet at intervals of length $\Delta > 0$, so that period k corresponds to $t \in [k\Delta, (k+1)\Delta]$. Play within this interval is as follows:

¹⁹By Lemma 13, this holds whenever $\delta \geq \bar{\delta}$ (in addition to Assumption 1).

- At time $k\Delta^+$, a public randomization device realizes, and the principal decides whether to delegate or freeze.
- Over the course of $(k\Delta, (k+1)\Delta)$, the agent sees the project arrival process. He then has in hand the best available project type:

$$\theta_k := \begin{cases} \bar{\theta} & \text{if at least one good project arrived during } (k\Delta, (k+1)\Delta), \\ \underline{\theta} & \text{otherwise.} \end{cases}$$

- If the principal delegated at time $k\Delta^+$, then the agent decides whether or not to adopt the project of type θ_k at time $(k+1)\Delta^-$. If the principal froze, the agent has no choice to make.

An initiated project of type θ provides the players a lump-sum revenue of θ , at a lump-sum cost (to the principal) of c .

Up to a strategically irrelevant scaling of payoffs, the above continuous time game is equivalent to one in discrete time with parameters $\delta_\Delta := e^{-r\Delta}$, $h_\Delta := 1 - e^{-\eta\Delta}$. Given fixed $\bar{\theta} > c > \underline{\theta} > 0$, we can invoke an appropriate lower bound²⁰ on $\frac{\eta}{r}$ to guarantee that Assumption 1 holds and that $\delta_\Delta \geq \bar{\delta}$ for sufficiently small Δ .

Fix a sequence of period lengths $\Delta \rightarrow 0$ such that the equilibrium value set converges.²¹ Let $\bar{v}_0$ be the highest agent value in the limit.

Adapting Theorem 2 tells us that $\bar{v}_0 > \omega_0$, where $\omega_0 := \eta(\bar{\theta} - \underline{\theta}) = \lim_{\Delta \rightarrow 0} \frac{\delta_\Delta}{1-\delta_\Delta} \omega_\Delta$, the **marginal value of (instantaneous) search**. Corollary 3 then tells us that the path of play in Pareto efficient equilibrium has several interpretable features in the limit. ‘No superfluous randomization’ tells us that, when the agent’s continuation value is at least ω_0 , no public randomization is currently used. ‘Good behavior with minimal punishment’ says that, if $v_t \in [\omega_0, \bar{v})$, then the principal delegates, the agent value law of motion is

$$\begin{cases} v \text{ follows } \dot{v}_t = r(v_t - \omega_0) & \text{while the agent doesn't adopt a project,} \\ v \text{ jumps to } v_t - r\underline{\theta} & \text{if the agent adopts a project,} \end{cases}$$

and the agent adopts only good projects. ‘Gift-giving as a last resort’ means that, when the

²⁰ $\frac{\eta}{r} > \frac{\underline{\theta}}{\bar{\theta} - \underline{\theta}}$ will imply Assumption 1. Given this, $\frac{\eta}{r}(\bar{\theta} - \underline{\theta})(1 - \frac{\varepsilon}{\bar{\theta}}) > c - \underline{\theta}$ will imply $\delta_\Delta \geq \bar{\delta}$.

²¹ Some such sequence exists, where convergence is taken with respect to the Hausdorff metric. Indeed, each $\mathcal{E}_\Delta^*$ belongs to the compact set $\{\mathcal{E} \subseteq \bar{\mathcal{E}} : \mathcal{E} \text{ nonempty closed convex}\}$, where $\bar{\mathcal{E}} := \{(v, b) \in \mathbb{R}_+^2 : v \leq \eta\bar{\theta} + b\underline{\theta}, \pi(v, b) \geq 0\}$.

continuation value hits $\bar{v}$, the principal continues to delegate to the agent, who immediately adopts a project (almost surely a bad one), and the agent's value jumps to $\bar{v} - r\theta$. Lastly, 'absorbing punishment' tells us that, once the agent's continuation value is below ω_0 , it stays there forever.

6.2 Defining the DCB contract

The DCB contract is characterized by a budget cap $\bar{x} \geq 0$ and initial budget balance $x \in [-1, \bar{x}]$, and consists of two regimes. At any time, players follow Controlled Budgeting or Capped Budgeting, depending on the agent's balance, x . The account balance can be understood as the number of projects the agent can initiate without immediately affecting the principal's delegation decisions.²²

Capped Budget ($x > 0$)

The account balance grows at the interest rate r as long as $x < \bar{x}$. Accrued interest is used to reward the agent for fiscal restraint. Since the opportunity cost of taking a project decreases in the account balance, the reward for diligence is increasing (exponentially) to maintain incentives. While there are funds in $\mathcal{A}$'s account, $\mathcal{P}$ fully delegates project choice to $\mathcal{A}$. However, every project that $\mathcal{A}$ initiates reduces the account balance to $x - 1$ (whether or not the latter is positive). Good projects being scarce, there are limits to how many projects the principal can credibly promise. When the balance is at the cap, the account can grow no further; accordingly, the agent takes a project immediately, yielding a balance of $\bar{x} - 1$.

Controlled Budget ($x \leq 0$)

The controlled budget regime is tailored to provide low revenue, low enough to be feasibly provided in aligned equilibrium. When $x < 0$, the agent is over budget, and the principal punishes the agent—more severely the further over budget the agent is—with a freeze, restoring the balance to zero. The continuation contract when the balance is $x = 0$ is some optimal aligned equilibrium.

Definition. *The Dynamic Capital Budgeting (DCB) contract $\sigma^{x, \bar{x}}$ is as follows:*

1. *The Capped Budget regime: $x > 0$.*

²²While the continuous time model eases exposition, it is not necessary for a DCB implementation. In the discrete time analogue, public randomization is used when the agent's balance lies between $\bar{x}$ and $(1 - \delta)\bar{x}$; the implementation is otherwise unchanged.

- While $x \in (0, \bar{x})$: $\mathcal{P}$ delegates, and $\mathcal{A}$ takes any available good projects and no bad ones. If $\mathcal{A}$ initiates a project, the balance jumps from x to $x - 1$; if $\mathcal{A}$ does not take a project, x drifts according to $\dot{x} = rx > 0$.
- When x hits $\bar{x}$: $\mathcal{P}$ delegates, and $\mathcal{A}$ takes a project immediately. If $\mathcal{A}$ adopts a project, the balance jumps from $\bar{x}$ to $\bar{x} - 1$; if $\mathcal{A}$ doesn't take a project, the balance remains at $\bar{x}$.

2. The **Controlled Budget** regime: $x \leq 0$.

- If $x \in [-1, 0)$: $\mathcal{P}$ freezes for duration $\frac{1}{r} \log \frac{\omega_0}{\omega_0 - \underline{\theta}|x|}$.
- After this initial freeze, $\mathcal{P}$ repeats the same policy forever: delegate until the next project, and then freeze for duration $\bar{\tau} = \frac{1}{r} \log \frac{\omega_0}{\omega_0 - \underline{\theta}}$.

In this regime, $\mathcal{A}$ adopts every good project and no bad projects.

Consider a sample path of the relationship between a hiring physics department and its university. At the department's inception, the university allocates a budget of three hires, with a cap of ten. Over time, the physics department searches for candidates. Every time the department finds an appropriate candidate, it hires—and the provost rubber stamps it—spending from the agreed-upon budget. Figure 4 represents one possible realized path of the account balance over time.

The department finds two suitable candidates to hire in its first year; some interest having accrued, the department budget is now at two hires. After the first year, the department enters a dry spell: finding no suitable candidate for six years, the department hires no one. Due to accrued interest, the account budget has increased dramatically from two to over eight hires. In its ninth year, the account hits the cap and can grow no further. The department can immediately hire up to nine physicists and continue to search (with its remaining budget) or it can hire ten candidates and enact a regime change by the provost. The department chooses to hire one physics professor (irrespective of quality) immediately, and continue to search with a balance of nine.

Over the next few years, the department is flush and hires many professors. First, for three years, the department hits its cap several times, hiring many mediocre candidates. After its eleventh year, the department faces a lucky streak, finding many great physicists over the following years, bringing the budget to one hire. In the next twelve years, the department finds few candidates worth hiring. However, the interest accrual is so slow that the physics department still depletes its budget, in the twenty-eighth year. Throughout this

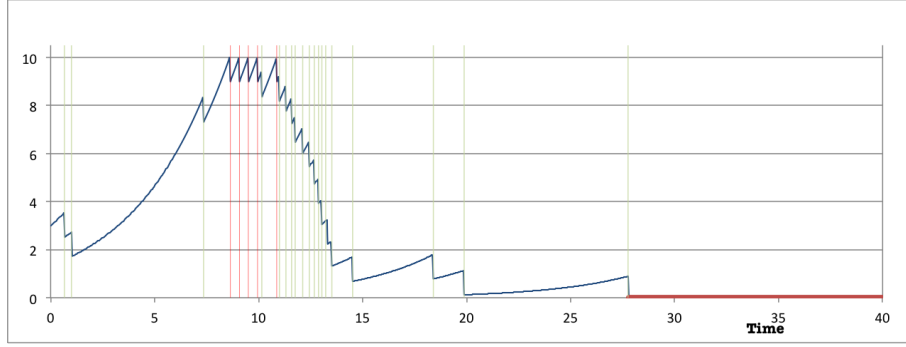


Figure 4: One realization of the balance's path under Controlled Budgeting (with $\bar{x} = 10$). Bad projects are clustered, and the account eventually runs dry.

initial phase, the department hires a total of twenty-four physics professors (many more than the account cap of ten).

At this point, the relationship changes permanently. After a temporary hiring freeze, the provost allows the department to resume its search, but follows any hire with a two-year hiring freeze. The relationship is now of a much more conservative character.

Notice that bad projects are clustered: the high balance $\bar{x} - 1$ just after a bad project means that the balance is likely to reach $\bar{x}$ again before the next arrival of a good project. So the next adopted project is likely bad. Given exponential growth, this effect is stronger the higher is the cap. In the Capped Budget regime, for a given account cap, the balance has non-monotonic profit implications. If the account runs low, there is an increased risk of imminently moving toward the low-revenue Controlled Budget regime. If the account runs high, the principal faces more bad projects in the near future. Controlled Budgeting is absorbing: once the balance falls low enough—which it eventually does—the agent will never take a bad project again.

Proposition 2. *Fixing an account cap and initial balance $\bar{x} > x > 0$, consider the Dynamic Capital Budget contract $\sigma^{x, \bar{x}}$.*

1. *Expected discounted revenue is $v(x) = \omega_0 + r\underline{\theta}x$.*
2. *Expected discounted number of bad projects is $b(x) = b^{\bar{x}}(x)$, uniquely determined by the delay differential equation*

$$(1 + \frac{\eta}{r})b(x) = \frac{\eta}{r}b(x - 1) + xb'(x),$$

with boundary conditions:

$$\begin{aligned} b|_{(-\infty, 0]} &= 0 \\ b(\bar{x}) - b(\bar{x} - 1) &= 1. \end{aligned}$$

3. $\sigma^{x, \bar{x}}$ is an equilibrium if and only if it exhibits nonnegative profit at the cap—that is,

$$\bar{\pi}(\bar{x}) := \pi(\omega_0 + r\theta\bar{x}, b(\bar{x})) \geq 0.$$

Proof. The first point follows from substituting into the v promise-keeping constraint, and verifying that (by direct computation) $\sigma^{0, \bar{x}}$ yields revenue ω_0 . The second point follows from Proposition 4 in Section 10.

For the third part, $v(x) - v(x - 1) = [\omega_0 + r\theta x] - [\omega_0 + r\theta(x - 1)] = r\theta$ at every x , so that the agent is always indifferent between taking or leaving a bad project. Thus, $\sigma^{x, \bar{x}}$ is an equilibrium if and only if it satisfies principal participation after every history. Revenue is linear, and b is (again by Proposition 4) convex. Therefore, profit is concave in x , in addition to being nonnegative at nonpositive balances. So, profit is nonnegative for all on-path balances if and only if it is nonnegative at the top. $\square$

6.3 Optimality of the DCB Contract

Finally we note that the DCB contract (with appropriate cap and initial balance) implements exactly the Pareto efficient equilibria. In this sense, our model gives an agency foundation to dynamic budgeting of capital expenditures. From the analysis of Subsection 6.1, the following is immediate.

Theorem 3. *Every Pareto efficient equilibrium payoff can be implemented as a DCB contract with cap $\bar{x} := \frac{\bar{v} - \omega_0}{r\theta}$ and some initial balance.*

Adapting Corollary 3 gives a partial uniqueness result as well: every Pareto efficient equilibrium follows a two-regime structure, with Capped Budgeting (the cap computed as above) as its first regime.

6.4 Comparative Statics

In light of Theorem 3, a principal-optimal contract is characterized by two elements: how much freedom the principal can credibly give the agent (the cap), and how much freedom the principal initially gives the agent (the initial balance). The first describes the equilibrium payoff set, while the second selects a principal-optimal equilibrium therein. In this subsection, we ask how these features change as the environment the players face varies. As parameters of the model change, and the pool of projects becomes more valuable, the agent enjoys greater sovereignty, with both the balance cap and the optimal initial balance increasing.

Proposition 3. *For any profile of parameters that satisfy our Assumptions (as in footnote 20), define the account cap $\bar{X}(\bar{\theta}, \underline{\theta}, c)$, and the initial account balance $X^*(\bar{\theta}, \underline{\theta}, c)$ employed for principal-optimal equilibrium. Both functions are increasing in the revenue parameters $\bar{\theta}, \underline{\theta}$, and decreasing in the cost parameter c .*

The proof is in Section 11 of the appendix. We first analyze a slight increase in a revenue parameter and observe that the profit at the original cap increases. This implies that a slight cap increase maintains the principal's credibility, and it is now an equilibrium. This delivers the first half of our comparative statics result: the new equilibrium has a higher account cap, offering greater flexibility to the agent. As the account cap increases, the frontier of the equilibrium set flattens at each account balance. Moreover, as revenue parameters increase, the principal's isoprofit curves can only get steeper. Accordingly, the unique tangency between the equilibrium frontier and an isoprofit occurs at a higher balance. This delivers the remaining comparative statics result: the agent is also given more initial leeway. A similar analysis applies to a cost reduction.

Of particular interest is the role of $\underline{\theta}$ in determining the optimal DCB account structure. On the one hand, the principal suffers less from a bad project when $\underline{\theta}$ is higher; on the other, the agent is more tempted. We show that the former effect always dominates in determining how much freedom the principal optimally gives the agent.

7 Extensions

In this section, we briefly describe some extensions to our model. The proofs are straightforward and omitted.

Monetary Transfers

We maintain the assumption of limited liability: $\mathcal{A}$ cannot give $\mathcal{P}$ money. If $\mathcal{P}$ can reward $\mathcal{A}$'s fiscal restraint through direct transfers, one of two things happens: (i) nothing changes and money is not used; or (ii) money simply replaces bad projects as a reward if money is more cost-effective. Which is more efficient depends on the relative size of the marginal cost $(c - \underline{\theta})B'(\bar{v})$ of allowing the agent to initiate bad projects and the marginal reward $\underline{\theta}$ of doing so. If providing monetary incentives is optimal, a modified DCB contract is used. The cap is raised to ensure zero profit with the new, more efficient incentivizing technology, and the agent is paid a flow of cash whenever his balance is at the cap. This modified DCB contract is reminiscent of the optimal contract in Biais et al. (2010).

Permanent Termination

In many applications, being in a given relationship automatically entails delegating. If a client hires a lawyer, she delegates the choice of hours to be worked. To stop delegating is to terminate the relationship, giving both players zero continuation value.²³ That is, at any moment, the principal must choose between fully delegating and ending the game forever. At first glance, this constraint may seem an additional burden on the principal. However, given Theorem 1 (together with the observation that agent payoffs below ω can be provided with stochastic termination), we see that it changes nothing. Indeed, replacing any temporary freeze—guaranteed to happen only in the absorbing “ $v \leq \omega$ ” regime—with stochastic termination leaves payoffs and incentives unchanged.

Agent Replacement

Suppose the principal has a means to punish the agent without punishing herself: the principal can fire the agent and immediately hire a new one. The credibility of the threat of replacement takes us far from our leading examples: for instance, the state government cannot sever its relationship with one of its counties.

A fired agent gets a continuation payoff of zero.²⁴ Every time the principal hires a new agent, she proposes a new contract. Any inefficiency that the principal faces, as well as any interesting relationship dynamics in the contracts, vanish: in any equilibrium, the principal

²³The agent could have a positive outside option. As long as it is below $\omega - (1 - \delta)\underline{\theta}$, the same argument holds.

²⁴Again, a positive agent outside option below $\omega - (1 - \delta)\underline{\theta}$ would change nothing.

always delegates to the current agent, who, in turn, exercises fiscal restraint. Indeed, the principal can simply fully delegate to the current agent and replace him as soon as he takes a project. By Assumption 1, the agent finds it incentive compatible to take only good projects. This policy, featuring a far less dynamic relationship than in our model, yields first-best profits for the principal. Adapting Theorem 1, one can show that every equilibrium of this modified model entails a value of at most ω for each agent.

Cheap Talk

What if we alter the order of play but preserve the core tension of our model? The agent privately observes the project type and sends the principal a message (e.g. a recommendation); the principal then chooses whether or not to fund the project.

First, consider the case where the principal can commit to an arbitrary stochastic implementation rule within-period, but cannot commit beyond a given period. Thus, the stage game is a full-fledged delegation problem, but the principal still cannot commit to long-run rewards. The results and insights are then qualitatively very similar, but not identical, to our model. Optimal play still comprises two regimes, the first entailing flexibility and the second entailing conservatism. The equation defining the equilibrium payoff frontier changes only for agent payoffs in $(\delta\bar{v} + (1 - \delta)\omega, \bar{v})$. At such payoffs, public randomization is replaced with stochastic implementation of bad projects. In particular, nothing at all changes in the limiting continuous time model of Section 6.

Now, what if the principal cannot commit to a project adoption choice and is therefore subject to interim incentive constraints? One can verify that the associated equilibrium payoff set lies somewhere between that of Theorem 2 and that of the one-period commitment case above. Therefore, again, this model's insights are qualitatively similar to those of our model, with the DCB contract remaining optimal in the limiting model of Section 6.

8 Final Remarks

In this paper, we have presented an infinitely repeated instance of the delegation problem. The agent will not represent the principal's interests without being offered dynamic incentives, while the principal cannot credibly commit to long-term rewards. Our principal's credibility concerns—an important ingredient of our leading applications—complicate the problem and yield substantively different economic outcomes.

First, we speak to the evolution of the relationship under Pareto efficient equilibrium. Early on, the relationship is highly productive but low-yield: the agent adopts every good project, but some bad projects as well. The lack of principal commitment limits the magnitude of credible promises, making this phase a transient one. As the relationship matures, it is high-yield but less productive: the agent adopts only good projects, but some good opportunities go unrealized. In this sense, the relationship transitions toward conservatism.

Second, we solve for the exact form of an efficient intertemporal delegation rule, our Dynamic Capital Budget contract, which comprises two regimes. In the first regime, Capped Budgeting, the agent has an expense account, which grows at the interest rate so long as its balance is below its cap; the principal fully delegates, with every project being financed from the account. The agent takes every available good project; only when at the cap does he adopt projects indiscriminately. Eventually, the account runs dry, and the players transition to the second regime, Controlled Budgeting, wherein they play a (Pareto inefficient) aligned equilibrium. Not only is the DCB contract profit-maximizing, but it in fact traces out the whole equilibrium value set;²⁵ we note that the analysis and results apply at any fixed discount rate.

While our main applications concern organizational economics outside of the firm, our results also speak to the canonical firm setting. There, the conflict of interest in our model could reflect an empire-building motive on the part of a department, or it may be an expression of the Baumol (1968) sales-maximization principle. If the relationship between a firm and one of its departments proceeds largely via delegation, then we shed light on the dynamic nature of this relationship. In doing so, we provide a foundation for dynamic budgeting within the firm.

²⁵More precisely, every equilibrium payoff is given by a mixture between the stage game equilibrium and a DCB contract.

Appendix

This appendix provides formal proof for the results on which the main text draws. First, we provide a characterization of the equilibrium payoff set of the discrete-time game. Next, we prove several useful properties of the Delay Differential Equation which characterizes the frontier of the limit equilibrium set. Finally, we derive comparative statics results for the optimal cap and initial balance of a Dynamic Capital Budget contract.

9 APPENDIX: Characterizing Equilibrium Values

In the current section, we characterize the equilibrium value set in our discrete time repeated game. As in the main text, we find it convenient to study payoffs in terms of *agent value* and *bad projects*. Accordingly, for any strategy profile σ , we let

$$\begin{aligned} v(\sigma) &= \mathbb{E}^\sigma \left[(1 - \delta) \sum_{k=0}^{\infty} \delta^k \mathbf{1}_{\{\text{a project is adopted in period } k\}} \theta_k \right]; \\ b(\sigma) &= \mathbb{E}^\sigma \left[(1 - \delta) \sum_{k=0}^{\infty} \delta^k \mathbf{1}_{\{\text{a project is adopted in period } k\}} \mathbf{1}_{\theta_k = \underline{\theta}} \right]. \end{aligned}$$

Below, we will analyze the perfect public equilibrium (PPE) value set,

$$\mathcal{E}^* = \{(v(\sigma), b(\sigma)) : \sigma \text{ is a PPE}\} \subseteq \mathbb{R}_+^2.$$

9.1 Self-Generation

To describe the equilibrium value set $\mathcal{E}^*$, we rely heavily on the machinery of Abreu, Pearce, and Stacchetti (1990), called APS hereafter. To provide the players a given value $y = (v, b)$ from today onward, we factor it into a (possibly random) choice of what happens today, and what the continuation will be starting tomorrow. What happens today depends on the probability (p) that the principal delegates, the probability ($\bar{a}$) of project adoption if a project is good, and the probability (a) of project adoption if a project is bad. The continuation values may vary based on what happens today: the principal may choose to freeze (y_F), the principal may delegate and agent may take a project (y_P), or the principal may delegate and agent may not take a project (y_N). Since the principal does not observe project types, these are the only three public outcomes.

We formalize this factorization in the following definition and theorem.

Definition 3. Given $Y \subseteq \mathbb{R}^2$:

- Say $y \in \mathbb{R}^2$ is **purely enforceable** w.r.t. Y if there exist $p, \bar{a}, a \in [0, 1]$ and $y_F, y_P, y_N \in Y$ such that:²⁶

1. (Promise keeping):

$$\begin{aligned} y &= (1-p)\delta y_F + ph \left\{ \bar{a} \left[(1-\delta)(\bar{\theta}, 0) + \delta y_P \right] + (1-\bar{a})\delta y_N \right\} \\ &\quad + p(1-h) \left\{ a \left[(1-\delta)(\underline{\theta}, 1) + \delta y_P \right] + (1-a)\delta y_N \right\} \\ &= (1-p)\delta y_F + p \left\{ h\bar{a} \left[(1-\delta)(\bar{\theta}, 0) + \delta(y_P - y_N) \right] \right. \\ &\quad \left. + (1-h)a \left[(1-\delta)(\underline{\theta}, 1) + \delta(y_P - y_N) \right] + \delta y_N \right\}. \end{aligned}$$

2. (Incentive-compatibility):

$$\begin{aligned} p &\in \arg \max_{\hat{p} \in [0,1]} (1-\hat{p})\delta\pi(y_F) + \hat{p} \left\{ h\bar{a} \left[(1-\delta)(\bar{\theta} - c) + \delta[\pi(y_P) - \pi(y_N)] \right] \right. \\ &\quad \left. + (1-h)a \left[(1-\delta)(\underline{\theta} - c) + \delta[\pi(y_P) - \pi(y_N)] \right] + \delta\pi(y_N) - \delta\pi(y_F) \right\}, \\ \bar{a} &\in \arg \max_{\hat{a} \in [0,1]} \hat{a} \left\{ (1-\delta)\bar{\theta} + \delta[v(y_P) - v(y_N)] \right\}, \\ a &\in \arg \max_{\hat{a} \in [0,1]} \hat{a} \left\{ (1-\delta)\underline{\theta} + \delta[v(y_P) - v(y_N)] \right\}. \end{aligned}$$

- Say $y \in \mathbb{R}^2$ is **enforceable** w.r.t. Y if there exists a Borel probability measure μ on $\mathbb{R}^2$ such that

$$1. y = \int_{\mathbb{R}^2} \hat{y} d\mu(\hat{y}).$$

2. $\hat{y}$ is purely enforceable almost surely with respect to $\mu(\hat{y})$.

- Let $W(Y) := \{y \in \mathbb{R}^2 : y \text{ is enforceable with respect to } Y\}$.
- Say $Y \subseteq \mathbb{R}^2$ is **self-generating** if $Y \subseteq W(Y)$.

Adapting methods from Abreu, Pearce, and Stacchetti (1990), one can readily characterize $\mathcal{E}^*$ via self-generation, through the following collection of results.

Lemma 2. Let W be as defined above.

- The set operator $W : 2^{\mathbb{R}^2} \longrightarrow 2^{\mathbb{R}^2}$ is monotone.
- $\mathcal{E}^*$ is the largest bounded self-generating set.

²⁶With a slight abuse of notation, for a given $y = (y_1, y_2) \in \mathbb{R}^2$, we will let $v(y) := y_1$ and $b(y) := y_2$.

- $W(\mathcal{E}^*) = \mathcal{E}^*$.
- Let $Y_0 \subseteq \mathbb{R}^2$ be any bounded superset of $\mathcal{E}^*$.²⁷ Define the sequence $(Y_n)_{n=1}^\infty$ recursively by $Y_n := W(Y_{n-1})$ for each $n \in \mathbb{N}$. Then $\bigcap_{n=1}^\infty Y_n = \mathcal{E}^*$.

9.2 A Cleaner Characterization

In light of the above, understanding the operator W will enable us to describe $\mathcal{E}^*$. That said, the definition of W is cumbersome. For the remainder of the current section, we work to better understand it.

Before doing anything else, we restrict attention to a useful domain for the map W .

Notation. Let $\mathcal{Y} := \{Y \subseteq \mathbb{R}_+^2 : \vec{0} \in Y, Y \text{ is compact and convex, and } \pi|_Y \geq 0\}$.

We need to work only with potential value sets in $\mathcal{Y}$. Indeed, the feasible set $\bar{\mathcal{E}}$ belongs to $\mathcal{Y}$, and it is straightforward to check that W takes elements of $\mathcal{Y}$ to $\mathcal{Y}$. Since $\mathcal{Y}$ is closed under intersections, we then know from the last bullet of Lemma 2, that $\mathcal{E}^* \in \mathcal{Y}$.

In seeking a better description of W , the following auxiliary definitions are useful.

Definition 4. Given $a \in [0, 1]$:

- Say $y \in \mathbb{R}_+^2$ is **a -Pareto enforceable** w.r.t. Y if there exist $y_P, y_N \in Y$ such that:

1. (Promise keeping):

$$y = h \left[(1 - \delta)(\bar{\theta}, 0) + \delta(y_P - y_N) \right] + (1 - h)a \left[(1 - \delta)(\underline{\theta}, 1) + \delta(y_P - y_N) \right] + \delta y_N.$$

2. (Agent incentive-compatibility):

$$\begin{aligned} 1 &\in \arg \max_{\hat{a} \in [0, 1]} \hat{a} \left\{ (1 - \delta)\bar{\theta} + \delta[v(y_P) - v(y_N)] \right\}, \\ a &\in \arg \max_{\hat{a} \in [0, 1]} \hat{a} \left\{ (1 - \delta)\underline{\theta} + \delta[v(y_P) - v(y_N)] \right\}. \end{aligned}$$

3. (Principal participation): $\pi(y) \geq 0$.

- Let $W_a(Y) := \{y \in \mathbb{R}_+^2 : y \text{ is } a\text{-Pareto enforceable w.r.t. } Y\}$.
- Let $W_f(Y) := \delta Y$.
- Let $\hat{W}(Y) := W_f(Y) \cup \bigcup_{a \in [0, 1]} W_a(Y)$. If Y is compact, then so is $\hat{W}(Y)$.²⁸

²⁷This can be ensured, for instance, by letting Y_0 contain the feasible set.

²⁸Indeed, it is the union of δY and a projection of the compact set $\{(a, y) \in [0, 1] \times \mathbb{R}^2 : y \text{ is } a\text{-Pareto enforceable w.r.t. } Y\}$.

The set $\hat{W}(Y)$ captures the enforceable (without public randomizations) values w.r.t. Y if:

1. The principal uses a pure strategy.
2. We relax principal IC to a participation constraint.
3. If the principal delegates and the project is good, then the agent takes the project.

The following proposition shows that, for the relevant $Y \in \mathcal{Y}$, it is without loss to focus on $co\hat{W}$ instead of W . The result is intuitive. The first two points are without loss because the principal's choices are observable. Toward (1), her private mixing can be replaced with public mixing with no effect on $\mathcal{A}$'s incentives. Toward (2), if the principal faces nonnegative profits with any pure action, she can be incentivized to take said action with stage Nash (min-max payoffs) continuation following the other choice. Toward (3), the agent's private mixing is not (given (2)) important for the principal's IC, and so we can replace it with public mixing between efficient (i.e. no good project being passed up) first-stage play and an initial freeze.

Lemma 3. *If $Y \in \mathcal{Y}$, then $W(Y) = co\hat{W}(Y)$.*

Proof. First, notice that $\delta Y \subseteq W(Y) \cap co\hat{W}(Y)$. It is a subset of the latter by construction, and of the former by choosing $y_P = y_N = \vec{0}$, $p = 0$, $\bar{a} = a = 1$, and letting y_F range over Y .

Take any $y \in \hat{W}(Y) \setminus \delta Y$. Then y is a -Pareto enforceable w.r.t. Y for some $a \in [0, 1]$, say witnessed by $y_P, y_N \in Y$. Letting $p = 1$, $\bar{a} = 1$, and $y_F = \vec{0} \in Y$, it is immediate that $p, \bar{a}, a \in [0, 1]$ and $y_F, y_P, y_N \in Y$ witness y being purely enforceable w.r.t. Y . Therefore, $y \in W(Y)$. So $\hat{W}(Y) \subseteq W(Y)$. The latter being convex, $co\hat{W}(Y) \subseteq W(Y)$ as well.

Take any extreme point y of $W(Y)$ which is not in δY . Then y must be purely enforceable w.r.t. Y , say witnessed by $p, \bar{a}, a \in [0, 1]$ and $y_F, y_P, y_N \in Y$. First, if $p\bar{a} = 0$, then²⁹

$$y = (1 - p)\delta y_F \in co\{\vec{0}, \delta y_F\} \subseteq \delta Y \subseteq \hat{W}(Y).$$

Now suppose $p\bar{a} > 0$, and define $a^* := \frac{a}{\bar{a}} \in [0, 1]$ and

$$y^* := h \left[(1 - \delta)(\vec{0}, 0) + \delta(y_P - y_N) \right] + (1 - h)a^* \left[(1 - \delta)(\vec{0}, 1) + \delta(y_P - y_N) \right] + \delta y_N.$$

Observe that y_P, y_N witness $y^* \in W_{a^*}(Y)$:

1. Promise keeping follows immediately from the definition of y^* .
2. Agent IC follows from agent IC in enforcement of y , and from the fact that incentive constraints are linear in action choices. As $\bar{a} > 0$ was optimal, $\bar{a}^* = 1$ is optimal here as well.

²⁹If $p\bar{a} = 0$, then either $p = 0$ or $\bar{a} = 0$. If $\bar{a} = 0$, then agent IC implies $a = 0$. So either $p = 0$ or $a = \bar{a} = 0$; in either case, promise keeping then implies $y = (1 - p)\delta y_F$.

3. Principal participation follows from principal IC in enforcement of y , and from the fact that $\pi(y_F) \geq 0$ because $\pi|_Y \geq 0$.

Therefore $y^* \in W_{a^*}(Y)$, from which it follows that

$$y = (1 - p)\delta y_F + p\bar{a}y^* \in co\{\delta y_F, y^*, \vec{0}\} \subseteq co\hat{W}(Y).$$

As every extreme point of $W(Y)$ belongs to $co\hat{W}(Y)$, and the latter is convex and (by Carathéodory's theorem) closed, all of $W(Y)$ belongs to $co\hat{W}(Y)$. $\square$

In view of the above proposition, we now only have to consider the much simpler map $\hat{W}$. As the following lemma shows, we can even further simplify, by showing that there is never a need to offer excessive punishment. That is, it is without loss to (1) make the agent's IC constraint (to resist bad projects) bind if he is being discerning, and (2) not respond to the agent's choice if he is being indiscriminate.

Lemma 4. Fix $a \in [0, 1]$, $Y \in \mathcal{Y}$, and $y \in \mathbb{R}^2$:

Suppose $a < 1$. Then $y \in W_a(Y)$ if and only if there exist $z_P, z_N \in Y$ such that:

1. (Promise keeping):

$$y = h[(1 - \delta)(\bar{\theta}, 0) + \delta(z_P - z_N)] + (1 - h)a[(1 - \delta)(\underline{\theta}, 1) + \delta(z_P - z_N)] + \delta z_N.$$

2. (Agent *exact* incentive-compatibility):

$$\delta[v(z_N) - v(z_P)] = (1 - \delta)\underline{\theta}.$$

3. (Principal participation): $\pi(y) \geq 0$.

Suppose $a = 1$. Then $y \in W_a(Y)$ if and only if there exists $z_N \in Y$ such that:

1. (Promise keeping):

$$y = (1 - \delta)[h(\bar{\theta}, 0) + (1 - h)(\underline{\theta}, 1)] + \delta z_N$$

2. (Principal participation): $\pi(y) \geq 0$.

Proof. In the first case, the “if” direction is immediate from the definition of W_a . In the second, it is immediate once we apply the definition of W_1 with $z_P = z_N$. Now we proceed to the “only if” direction.

Consider any $y \in W_a(Y)$, with y_P, y_N witnessing a -Pareto enforceability. Define

$$\bar{y} := [h + a(1 - h)]y_P + (1 - h)(1 - a)y_N \in Y.$$

So $\bar{y}$ is the on-path expected continuation value.

In the case of $a < 1$, define

$$\begin{aligned} q &:= \frac{(1 - \delta)\underline{\theta}}{\delta[v(y_N) - v(y_P)]} \quad (\in [0, 1], \text{ by IC}) \\ z_P &:= (1 - q)\bar{y} + qy_P \\ z_N &:= (1 - q)\bar{y} + qy_N. \end{aligned}$$

By construction, $\delta[v(z_N) - v(z_P)] = (1 - \delta)\underline{\theta}$, as desired.³⁰

In the case of $a = 1$, let $z_N := \bar{y}$ and $z_P := z_N$.

Notice that z_P, z_N witness $y \in W_a(Y)$. Promise keeping comes from the definition of $\bar{y}$, principal participation comes from the hypothesis that $W_a(Y) \ni y$, and IC (exact in the case of $a < 1$) comes by construction. $\square$

In the first part of the lemma, $\delta[v(y_N) - v(y_P)] \in (1 - \delta)[\underline{\theta}, \bar{\theta}]$ has been replaced with $\delta[v(z_N) - v(z_P)] = (1 - \delta)\underline{\theta}$. That is, it is without loss to make the agent's relevant incentive constraint—to avoid taking bad projects—bind. This follows from the fact that $Y \supseteq \text{co}\{y_P, y_N\}$. The second part of the lemma says that, if the agent is not being at all discerning, nothing is gained from disciplining him.

The above lemma has a clear interpretation: without loss of generality, the principal uses the minimal possible punishment. The lemma also yields the following:

Lemma 5. *Suppose $a \in (0, 1)$, $Y \in \mathcal{Y}$, and $y \in W_a(Y)$. Then there is some $y^* \in W_0$ such that*

$$v(y^*) = v(y) \text{ and } b(y^*) < b(y).$$

That is, $y_1^* = y_1$ and $y_2^* < y_2$.

Proof. Appealing to Lemma 4, there exist $z_P, z_N \in Y$ such that:

1. (Promise keeping):

$$y = h \left[(1 - \delta)(\bar{\theta}, 0) + \delta(z_P - z_N) \right] + (1 - h)a \left[(1 - \delta)(\underline{\theta}, 1) + \delta(z_P - z_N) \right] + \delta z_N$$

2. (Agent exact incentive-compatibility):

$$\delta[v(z_N) - v(z_P)] = (1 - \delta)\underline{\theta}$$

³⁰In the case of $a \in (0, 1)$, $q = 1$ (by agent IC), so that $z_P = y_P$ and $z_N = y_N$. The real work was needed for the case of $a = 0$.

3. (Principal participation): $\pi(y) \geq 0$.

Given agent exact IC, we know $v(z_N) > v(z_P)$. Let $z_P^* := \left(v(z_P), \min \left\{ b(z_P), \frac{v(z_P)}{v(z_N)} b(z_N) \right\} \right)$. As either $z_P^* = z_P$ or $z_P^* \in \text{co}\{\vec{0}, z_N\}$, we have $z^* \in Y$.

Let $y^* := h \left[(1 - \delta)(\bar{\theta}, 0) + \delta(z_P^* - z_N) \right] + \delta z_N$. Then

$$\begin{aligned} v(y) - v(y^*) &= (1 - h)a \left\{ (1 - \delta)\underline{\theta} + \delta[v(z_P^*) - v(z_N)] \right\} - h\delta[v(z_P^*) - v(z_P)] \\ &= (1 - h)a \left\{ (1 - \delta)\underline{\theta} + \delta[v(z_P) - v(z_N)] \right\} - h\delta 0 = 0, \end{aligned}$$

while

$$\begin{aligned} b(y) - b(y^*) &= (1 - h)a \left\{ (1 - \delta) + \delta[b(z_P) - b(z_N)] \right\} - h\delta[b(z_P^*) - b(z_P)] \\ &= (1 - h)a \left\{ (1 - \delta) + \delta[b(z_P^*) - b(z_N)] + \delta[b(z_P) - b(z_P^*)] \right\} - h\delta[b(z_P^*) - b(z_P)] \\ &= (1 - h)a \left\{ (1 - \delta) + \delta[b(z_P^*) - b(z_N)] \right\} + [h + (1 - h)a]\delta[b(z_P) - b(z_P^*)] \\ &\geq (1 - h)a > 0. \end{aligned}$$

Now, notice that z_P^*, z_N witness $y^* \in W_0(Y)$. Promise keeping holds by fiat, agent IC holds because $v(z_P^*) = v(z_P)$ by construction, and principal participation follows from

$$\pi(y^*) - \pi(y) = -\pi(0, b(y) - b(y^*)) > 0.$$

□

The above lemma is a strong bang-bang result. It is not simply sufficient to restrict attention to equilibria with no private mixing; it is necessary too. Any equilibrium in which the agent mixes on-path is Pareto dominated.

9.3 Self-Generation for Frontiers

Through Lemmata 3, 4, and 5, we simplified analysis of the APS operator W applied to the relevant value sets. In the current subsection, we profit from that simplification in characterizing the efficient frontier of $\mathcal{E}^*$. First, we make a small investment in some new notation.

Notation. Let $\mathcal{B}$ denote the space of functions $B : \mathbb{R} \rightarrow \mathbb{R}_+ \cup \{\infty\}$ such that: (1) B is convex, (2) $B(0) = 0$, (3) B 's proper domain $\text{dom}(B) := B^{-1}(\mathbb{R}_+)$ is a compact subset of $\mathbb{R}_+$, (4) B is continuous on $\text{dom}(B)$, and (5) $\pi(v, B(v)) \geq 0$ for every $v \in \text{dom}(B)$.

Just as $\mathcal{Y}$ is the relevant space of value sets, $\mathcal{B}$ is the relevant space of frontiers of value sets.

Notation. For each $Y \in \mathcal{Y}$, define the *efficient frontier function* of Y :

$$\begin{aligned} B_Y : \mathbb{R} &\longrightarrow \mathbb{R}_+ \cup \{\infty\} \\ v &\longmapsto \min\{b \in \mathbb{R}_+ : (v, b) \in Y\} \end{aligned}$$

It is immediate that for $Y \in \mathcal{Y}$, the function B_Y belongs to $\mathcal{B}$.

Notation. Define the following functions:

$$\begin{aligned} T : \mathcal{B} &\longrightarrow \mathcal{B} \\ \hat{B} &\longmapsto B_{W(\text{co}[\text{graph}(\hat{B})])} = B_{\text{co}\hat{W}(\text{co}[\text{graph}(\hat{B})])}, \\ T_f : \mathcal{B} &\longrightarrow \mathcal{B} \\ \hat{B} &\longmapsto B_{W_f(\text{co}[\text{graph}(\hat{B})])} = B_{\delta(\text{co}[\text{graph}(\hat{B})])}, \\ \text{and for } 0 \leq a \leq 1, \quad T_a : \mathcal{B} &\longrightarrow \mathcal{B} \\ \hat{B} &\longmapsto B_{W_a(\text{co}[\text{graph}(\hat{B})])}. \end{aligned}$$

These objects are not new. The map T [resp. T_f , T_a] is just a repackaging of W [resp. W_f , W_a], made to operate on frontiers of value sets, rather than on value sets themselves.

Lemmata 3, 4, and 5 help simplify our analysis of T , which in turn helps us characterize the efficient frontier of $\mathcal{E}^*$. We now proceed along these lines.

The following lemma is immediate from the definition of the map $Y \mapsto B_Y$.

Lemma 6. If $\{Y_i\}_{i \in \mathbb{I}} \subseteq \mathcal{Y}$, then $B_{\overline{\text{co}}[\bigcup_{i \in \mathbb{I}} Y_i]}$ is the convex lower envelope of $\inf_{i \in \mathbb{I}} B_{Y_i}$.³¹

The following proposition is the heart of our main characterization of the set $\mathcal{E}^*$'s frontier. It amounts to a complete description of the behavior of T .

Lemma 7. Fix any $B \in \mathcal{B}$ and $v \in \mathbb{R}$. Then:

1. $TB = \text{cvx} \left[\min \{T_f B, T_0 B, T_1 B\} \right]$.
2. For $i \in \{f, 0, 1\}$,

$$T_i B(v) = \begin{cases} \check{T}_i B(v) & \text{if } \pi(v, \check{T}_i^\Delta B(v)) \geq 0 \\ \infty & \text{otherwise,} \end{cases}$$

³¹The **convex lower envelope** of a function F is $\text{cvx}F$, the largest convex upper-semicontinuous function below it. Equivalently, $\text{cvx}F$ is the pointwise supremum of all affine functions below F .

where

$$\begin{aligned}
\check{T}_f B(v) &:= \delta B\left(\frac{v}{\delta}\right), \\
\check{T}_0 B(v) &:= \delta \left[hB\left(\frac{v - (1 - \delta)\theta_E}{\delta}\right) + (1 - h)B\left(\frac{v - (1 - \delta)[\theta_E - \underline{\theta}]}{\delta}\right) \right] \\
&= \delta \left[hB\left(\frac{v - (1 - \delta)\omega - (1 - \delta)\underline{\theta}}{\delta}\right) + (1 - h)B\left(\frac{v - (1 - \delta)\omega}{\delta}\right) \right], \\
\check{T}_1 B(v) &:= (1 - h)(1 - \delta) + \delta B\left(\frac{v - (1 - \delta)\theta_E}{\delta}\right) \\
&= (1 - h)(1 - \delta) + \delta B\left(\frac{v - (1 - \delta)\omega - (1 - \delta)\underline{\theta}}{\delta}\right).
\end{aligned}$$

Proof. That $TB = \text{cvx} \left[\min \{T_f B, \inf_{a \in [0,1]} T_a B\} \right]$ is a direct application of Lemma 6. Then, appealing to Lemma 5, $T_a B \geq T_0 B$ for every $a \in (0, 1)$. This proves the first point.

In what follows, let $Y := \text{co}[\text{graph}(B)]$ so that $TB = B_{W(Y)}$.

- Consider any $y \in W_0(Y)$:

Lemma 4 delivers $z_P, z_N \in Y$ such that

$$\begin{aligned}
y &= h \left[(1 - \delta)(\bar{\theta}, 0) + \delta(z_P - z_N) \right] + \delta z_N, \\
(1 - \delta)\underline{\theta} &= \delta[v(z_N) - v(z_P)].
\end{aligned}$$

Rewriting with $z_P = (v_P, b_P)$ and $z_N = (v_N, b_N)$, and rearranging yields:

$$\begin{aligned}
(1 - \delta)\underline{\theta} &= \delta[v_N - v_P] \\
(v, b) &= h \left[(1 - \delta)(\bar{\theta}, 0) + \delta(v_P - v_N, b_P - b_N) \right] + \delta(v_N, b_N) \\
&= h \left((1 - \delta)(\bar{\theta} - \underline{\theta}), \delta[b_P - b_N] \right) + \delta(v_N, b_N) \\
&= (\theta_E - \underline{\theta} + \delta v_N, h\delta b_P + (1 - h)\delta b_N).
\end{aligned}$$

Solving for the agent values yields

$$v_N = \frac{v - (1 - \delta)[\theta_E - \underline{\theta}]}{\delta} \text{ and } v_P = v_N - \delta^{-1}(1 - \delta)\underline{\theta} = \frac{v - (1 - \delta)\theta_E}{\delta}.$$

So given any $v \in \mathbb{R}_+$:

$$\begin{aligned}
T_0 B(v) &= \inf_{b, b_P, b_N} b \\
\text{s.t. } &\pi(v, b) \geq 0, \quad b = \delta [hb_P + (1-h)b_N], \\
&\text{and } \left(\frac{v - (1-\delta)[\theta_E - \underline{\theta}]}{\delta}, b_N \right), \left(\frac{v - (1-\delta)\theta_E}{\delta}, b_P \right) \in Y \\
&= \inf_{b, b_P, b_N} b = \delta [hb_P + (1-h)b_N] \\
\text{s.t. } &\pi(v, b) \geq 0 \text{ and } \left(\frac{v - (1-\delta)[\theta_E - \underline{\theta}]}{\delta}, b_N \right), \left(\frac{v - (1-\delta)\theta_E}{\delta}, b_P \right) \in Y \\
&= \begin{cases} b = \delta \left[hB \left(\frac{v - (1-\delta)\theta_E}{\delta} \right) + (1-h)B \left(\frac{v - (1-\delta)[\theta_E - \underline{\theta}]}{\delta} \right) \right] & \text{if } \pi(v, b) \geq 0, \\ \infty & \text{otherwise.} \end{cases}
\end{aligned}$$

- Consider any $y \in W_1(Y)$:

Lemma 4 now delivers $z_N = (v_N, b_N) \in Y$ such that

$$y = (1-\delta)[h(\bar{\theta}, 0) + (1-h)(\underline{\theta}, 1)] + \delta z_N,$$

which can be rearranged to

$$(v, b) = ((1-\delta)\theta_E + \delta v_N, (1-h)(1-\delta) + \delta b_N).$$

So given any $v \in \mathbb{R}_+$:

$$\begin{aligned}
T_1 B(v) &= \inf_{b, b_N} b \\
\text{s.t. } &\pi(v, b) \geq 0, \quad b = (1-h)(1-\delta) + \delta b_N, \text{ and } \left(\frac{v - (1-\delta)\theta_E}{\delta}, b_N \right) \in Y \\
&= \inf_{b_P, b_N} b = (1-h)(1-\delta) + \delta b_N \\
\text{s.t. } &\pi(v, b) \geq 0 \text{ and } \left(\frac{v - (1-\delta)\theta_E}{\delta}, b_N \right) \in Y \\
&= \begin{cases} b = (1-h)(1-\delta) + \delta B \left(\frac{v - (1-\delta)\theta_E}{\delta} \right) & \text{if } \pi(v, b) \geq 0, \\ \infty & \text{otherwise.} \end{cases}
\end{aligned}$$

- Lastly, given any $v \in \mathbb{R}_+$:

$$\begin{aligned}
T_f B(v) &= \inf_{b, b_N} b \\
\text{s.t. } &\pi(v, b) \geq 0, \quad b = \delta b_N, \text{ and } \left(\frac{v}{\delta}, b_N\right) \in Y \\
&= \begin{cases} b = \delta B\left(\frac{v}{\delta}\right) & \text{if } \pi(v, b) \geq 0, \\ \infty & \text{otherwise.} \end{cases}
\end{aligned}$$

□

9.4 The Efficient Frontier

In this subsection, we characterize the frontier $B_{\mathcal{E}^*}$ of the equilibrium value set. We first translate APS's self-generation to the setting of frontiers. This, along with Lemma 7 delivers our Bellman equation, Corollary 5. Then, in Lemma 9, we characterize aligned optimal equilibrium. Finally, in Lemma 11, we fully characterize the frontier $B_{\mathcal{E}^*}$.

Lemma 8. *Suppose $Y \in \mathcal{Y}$ with $W(Y) = Y$. Then $T B_Y = B_Y$.*

Proof. First, because W is monotone, $W(\text{co}[\text{graph}(B_Y)]) \subseteq W(Y) = Y$. Thus the efficient frontier of the former is higher than that of the latter. That is, $T B_Y \geq B_Y$.

Now take any $v \in \text{dom}(B_Y)$ such that $y := (v, B_Y(v))$ is an extreme point of Y . We want to show that $T B_Y(v) \leq B_Y(v)$.

By Lemma 3, $y \in W_f(Y) \cup \bigcup_{a \in [0,1]} W_a(Y)$.

- If $y \in W_f(Y)$, then $\frac{v}{\delta}, \vec{0} \in Y$, so that the extreme point y must be equal to $\vec{0}$. But in this case, $T B_Y(v) = T B_Y(0) = 0 = B_Y(0) = B_Y(v)$.
- If $y \in W_a(Y)$ for some $a \in [0, 1]$, say witnessed by $y_P, y_N \in Y$, then let

$$\begin{aligned}
z_P &:= (v(y_P), B_Y(v(y_P))) \\
z_N &:= (v(y_N), B_Y(v(y_N))) \\
z &:= h \left[(1 - \delta)(\vec{0}, 0) + \delta(z_P - z_N) \right] + (1 - h)a \left[(1 - \delta)(\underline{\theta}, 1) + \delta(z_P - z_N) \right] + \delta z_N.
\end{aligned}$$

Then

$$\begin{aligned}
b(z) &= (1 - h)(1 - \delta)a + [h + (1 - h)a]\delta B_Y(v(y_P)) + (1 - h)(1 - a)\delta B_Y(v(y_N)) \\
&\leq (1 - h)(1 - \delta)a + [h + (1 - h)a]\delta b(y_P) + (1 - h)(1 - a)\delta b(y_N) \\
&= b(y) = B_Y(v),
\end{aligned}$$

and z_P, z_N witness $z \in W_a(\text{co}[\text{graph}(B_Y)])$. In particular, $TB_Y(v) = TB_Y(v(z)) \leq b(z) \leq B_Y(v)$.

Next, consider any $v \in \text{dom}(B_Y)$. There is some probability measure μ on the extreme points of Y such that $(v, B_Y(v)) = \int_Y y \, d\mu(y)$. By minimality of $B_Y(v)$, it must be that $y \in \text{graph}(B_Y)$ a.s.- $\mu(y)$. So letting μ_1 be the marginal of μ on the first coordinate, $(v, B_Y(v)) = \int_{v(Y)} (u, B_Y(u)) \, d\mu_1(u)$, so that

$$B_Y(v) = \int_{v(Y)} B_Y \, d\mu_1 \geq \int_{v(Y)} TB_Y \, d\mu_1 \geq TB_Y(v),$$

where the last inequality follows from Jensen's inequality.

This completes the proof. □

The Bellman equation follows immediately.

Corollary 5. $B := B_{\mathcal{E}^*}$ solves the Bellman equation $B = \text{cvx} \left[\min \{T_f B, T_0 B, T_1 B\} \right]$.

Aligned Optimal Equilibrium In line with the main text, we now proceed to characterize the payoffs attainable in equilibria with no bad projects.

Lemma 9.

1. *There exist productive aligned equilibria if and only if Assumption 1 holds.*
2. *Every aligned equilibrium generates revenue of at most ω .*
3. *If Assumption 1 holds, then $(\omega, 0) \in \mathcal{E}^*$.*

Proof. We proceed in reverse, invoking Lemma 7 throughout.

For (3), define $B \in \mathcal{B}$ via $B(v) = 0$ for $v \in [0, \omega]$ and $B(v) = \infty$ otherwise. Notice that, $TB(\omega) \leq T_0 B(\omega) = \delta h B(\frac{\omega - (1-\delta)\theta_E}{\delta}) + \delta(1-h)B(\frac{\omega - (1-\delta)\omega}{\delta}) = 0$, where the second equality holds by Assumption 1. Because TB is convex, B is self-generating, and thus $B \geq B_{\mathcal{E}^*}$. (3) follows.

We now proceed to verify (2), i.e. that $\hat{v} > \omega$ implies that $(0, \hat{v}) \notin \mathcal{E}^*$.

Suppose $v > \omega$ has $B_{\mathcal{E}^*}(v) = 0$. Then $B_{\mathcal{E}^*}|_{[0,v]} = 0$, and

$$0 = B_{\mathcal{E}^*}(v) = TB_{\mathcal{E}^*}(v) = \min\{T_f B_{\mathcal{E}^*}(v), T_0 B_{\mathcal{E}^*}(v), T_1 B_{\mathcal{E}^*}(v)\}.$$

Notice that $TB_{\mathcal{E}^*}(v) \neq T_1 B_{\mathcal{E}^*}(v)$ as the latter is strictly positive. If $TB_{\mathcal{E}^*}(v) = T_f B_{\mathcal{E}^*}(v)$, then since $B_{\mathcal{E}^*}$ is increasing,

$$B_{\mathcal{E}^*} \left(v + \frac{1-\delta}{\delta}(v-\omega) \right) \leq B_{\mathcal{E}^*} \left(v + \frac{1-\delta}{\delta}v \right) = \delta^{-1} T_f B_{\mathcal{E}^*}(v) = 0.$$

Finally, if $TB_{\mathcal{E}^*}(v) = T_0B_{\mathcal{E}^*}(v)$, then³²

$$0 = B_{\mathcal{E}^*}\left(\frac{v - (1 - \delta)\omega}{\delta}\right) = \delta B_{\mathcal{E}^*}\left(v + \frac{1 - \delta}{\delta}(v - \omega)\right).$$

So either way, $B_{\mathcal{E}^*}\left(v + \frac{1 - \delta}{\delta}(v - \omega)\right) = 0$ too.

Now, if $\hat{v} > \omega$ with $(0, \hat{v}) \in \mathcal{E}^*$, then applying the above inductively yields a sequence $v_n \rightarrow \infty$ on which $B_{\mathcal{E}^*}$ takes value zero.³³ This would contradict the compactness of $B_{\mathcal{E}^*}$'s proper domain.

By (3) we know that productive aligned equilibria exist if Assumption 1 holds. To see the necessity of the assumption, suppose for a contradiction that it fails and yet some productive aligned equilibrium exists. Let v be the agent's continuation value at some on-path history at which the principal delegates. By (2), we know $v \leq \omega$. As Assumption 1 fails, we then know $\delta v < (1 - \delta)\theta$, contradicting agent IC. $\square$

Proof of Proposition 1

Proof. Given Assumption 1, the proposition follows directly from Lemma 9 above. $\square$

Optimal Equilibrium We now focus on the frontier of the whole equilibrium set. Before proceeding to the full characterization, we establish a single crossing result: indiscriminate project adoption is initially used only for the highest agent values.

Lemma 10. Fix $B \in \mathcal{B}$, and suppose $B^{-1}(0) = [0, \omega]$:

1. If $v > \omega$, then $T_0B(v) < T_fB(v)$ (unless both are ∞).
2. There is a cutoff $\underline{v} \geq \omega$ such that

$$\begin{cases} T_0B(v) \leq T_1B(v) & \text{if } v \in [\omega, \underline{v}); \\ T_0B(v) \geq T_1B(v) & \text{if } v > \underline{v}. \end{cases}$$

Proof. $B(\omega) = 0$, and B is convex. Therefore, B is strictly increasing above ω on its domain, so that³⁴ $\check{T}_fB(v) > \check{T}_0B(v)$, confirming the first point.

Given v ,

$$\frac{v - (1 - \delta)[\theta_E - \underline{\theta}]}{\delta} - \frac{v - (1 - \delta)\theta_E}{\delta} = \frac{1 - \delta}{\delta} \underline{\theta}$$

is a nonnegative constant.

³²Since a weighted average of two nonnegative numbers can be zero only if both numbers are zero.

³³Let $v_0 = \hat{v}$ and $v_{n+1} = v_n + \frac{1 - \delta}{\delta}(v_n - \omega) \geq \hat{v} + n(\hat{v} - \omega)$.

³⁴The relationship is as shown if $\check{T}_0B(v) < \infty$. Otherwise, $T_fB(v) = T_0B(v) = \infty$.

Since B is convex, it must be that the continuous function

$$v \mapsto \check{T}_0 B(v) - \check{T}_1 B(v) = \delta(1-h) \left[B\left(\frac{v - (1-\delta)[\theta_E - \underline{\theta}]}{\delta}\right) - B\left(\frac{v - (1-\delta)\theta_E}{\delta}\right) \right] - (1-h)$$

is increasing on its proper domain. The second point follows.³⁵

□

Our characterization of the equilibrium frontier B can now be stated.

Lemma 11 (Equilibrium Frontier).

Suppose Assumption 1 holds, and let $B := B_{\mathcal{E}^*}$ and $\bar{v} := \max \text{dom}(B)$.

1. $\bar{v} \geq \omega$, and $B(v) = 0$ for $v \in [0, \omega]$.

2. If $\bar{v} > \omega$, then

$$\begin{aligned} B(v) &= T_0 B(v) \text{ for } v \in [\omega, \delta\bar{v} + (1-\delta)(\theta_E - \underline{\theta})]; \\ B(v) &\text{ is affine in } v \text{ for } v \in [\delta\bar{v} + (1-\delta)(\theta_E - \underline{\theta}), \bar{v}]; \\ B(\bar{v}) &= T_1 B(\bar{v}). \end{aligned}$$

3. If $\bar{v} > \omega$, then $\pi(\bar{v}, B(\bar{v})) = 0$.

Proof. The first point follows directly from Lemma 9. Now suppose $\bar{v} > \omega$.

Let $\underline{v} := \delta\bar{v} + (1-\delta)(\theta_E - \underline{\theta})$. Any $v > \underline{v}$ has $\frac{v - (1-\delta)[\theta_E - \underline{\theta}]}{\delta} > \frac{\underline{v} - (1-\delta)[\theta_E - \underline{\theta}]}{\delta} = \bar{v}$, so that $B\left(\frac{v - (1-\delta)[\theta_E - \underline{\theta}]}{\delta}\right) = \infty$, and therefore (appealing to Lemma 7) $T_0 B(v) = \infty$. Therefore, the cutoff defined in Lemma 10 is $\leq \underline{v}$.

Since $T_0 B, T_1 B$ are both convex, there exist some $v_0, v_1 \in [\omega, \bar{v}]$ such that $v_0 \leq v_1, \underline{v}$, and:

$$\begin{aligned} B(v) &= 0 \text{ for } v \in [0, \omega]; \\ B(v) &= T_0 B(v) \text{ for } v \in [\omega, v_0]; \\ B(v) &\text{ is affine in } v \text{ for } v \in [v_0, v_1]; \\ B(v) &= T_1 B(v) \text{ for } v \in [v_1, \bar{v}]. \end{aligned}$$

Let $m > 0$ denote the left-sided derivative of B at v_1 (which is simply the slope of B on (v_0, v_1) if $v_0 \neq v_1$).

Let $[v_0, v_1]$ be maximal (w.r.t. set inclusion) such that the above decomposition is still correct.

Notice then that $v_1 = \bar{v}$. Indeed, for $v \in [v_1, \bar{v})$, the right-hand side derivative satisfies $m \leq (TB)'(v) = (T_1 B)'(v) = B\left(\frac{v - \theta_E}{\delta}\right)$, by convexity of B and Lemma 7. Minimality of v_0 then implies

³⁵Because wherever $\check{T}_i B(v) \geq \check{T}_j B(v)$, we have $T_i B(v) \geq T_j B(v)$ as well.

$\frac{v-\theta_E}{\delta} \geq v_0$, so that $\frac{v-\theta_E}{\delta} \in [v_0, v_1]$. Then, for v in a neighborhood of v_1 in $[v_1, \bar{v})$, we have $B'(v) = (T_1 B)'(v) = m$. Maximality of v_1 tells us that $v_1 = \bar{v}$.

Finally, we need to show that $v_0 = \underline{v}$. Now, by minimality of v_0 , it must be that for any $v \in [0, v_0)$, the right-side derivative $B'(v) < m$. If $v_0 < \underline{v}$ (so that $\frac{v_0-(1-\delta)[\theta_E-\underline{\theta}]}{\delta} < \underline{v}$), then Lemma 7 gives us

$$\begin{aligned} m &= B'(v_0) \leq T_0 B'(v_0) \\ &= hB'\left(\frac{v_0-(1-\delta)\theta_E}{\delta}\right) + (1-h)B'\left(\frac{v_0-(1-\delta)[\theta_E-\underline{\theta}]}{\delta}\right) \\ &= hB'\left(\frac{v_0-(1-\delta)\theta_E}{\delta}\right) + (1-h)m < m, \end{aligned}$$

a contradiction. Therefore $v_0 = \underline{v}$, and the second point of the theorem follows.

For the last point, assume $\bar{v} > \omega$ and yields strictly positive profits. This implies that $\pi(v, T_1 B(\bar{v})) = \pi(v, B(\bar{v})) > 0$. Then, for sufficiently small $\gamma > 0$, the function $B^\gamma \in \mathcal{B}$ given by $B^\gamma(v) = \begin{cases} B(v) & \text{if } v \in [0, \bar{v}] \\ T_1 B(v) & \text{if } v \in [\bar{v}, \bar{v} + \gamma] \end{cases}$ is self-generating, contradicting the fact that $\mathcal{E}^*$ is the largest self-generating set. $\square$

To better understand efficient equilibrium behavior, the following lemma is useful.

Lemma 12. *Let $\bar{v}, B$ be as in Lemma 11. If $\bar{v} > \omega$, then there are constants $\alpha, \beta > 0$ such that $B(\omega + \epsilon) = \alpha \epsilon^{1+\beta}$ for sufficiently small $\epsilon > 0$.*

Proof. For $\epsilon \in (0, \bar{v} - \omega)$, define $Q(\epsilon) := \frac{B(\omega+\epsilon)-B(\omega)}{\epsilon} = \frac{1}{\epsilon} B(\omega + \epsilon)$. For small enough $\epsilon > 0$,

$$Q(\epsilon) = \frac{\delta}{\epsilon} \left[(1-h)B\left(\frac{(\omega+\epsilon)-(1-\delta)\omega}{\delta}\right) + h0 \right] = (1-h)Q\left(\frac{\epsilon}{\delta}\right).$$

For $t \in (-\infty, \log(\bar{v} - \omega))$, let $q(t) := \log Q(e^t)$. Then q is a continuous function, and $q(t - \log \delta) = q(t) - \log(1-h)$ for low enough $t \in \mathbb{R}$. Therefore, far enough to the left of its domain, q is affine with slope $\beta := \frac{\log(1-h)}{\log \delta} > 0$. That is, there is a constant $\alpha > 0$ such that $q(t) = \log \alpha + \beta t$ for low enough t . Therefore, for small enough $\epsilon > 0$,

$$Q(\epsilon) = e^{q(\log \epsilon)} = e^{\log \alpha + \beta \log \epsilon} = \alpha \epsilon^\beta \implies B(\omega + \epsilon) = \alpha \epsilon^{1+\beta}.$$

Lastly, nonnegativity of B and Proposition 1 tell us $\alpha > 0$. $\square$

Lastly, the next lemma records a sufficient condition for existence of some equilibrium with an initial bad project.

Lemma 13. *Let $\bar{v}, B$ be as defined in Lemma 11. If Assumption 1 holds and $\delta \geq \bar{\delta}$, then $\bar{v} > \omega$.*

Proof. By Lemma 9, we know $B(\omega) = 0$, implying (in the notation of Lemma 11)

$$\check{T}_1 B((1 - \delta)\theta_E + \delta\omega) = (1 - \delta)(1 - h) + B(\omega) = (1 - \delta)(1 - h).$$

That $\delta \geq \bar{\delta}$ then tells us, by direct computation, that $\pi((1 - \delta)\theta_E + \delta\omega, (1 - \delta)(1 - h)) \geq 0$. Therefore, $T_1 B((1 - \delta)\theta_E + \delta\omega) = \check{T}_1 B((1 - \delta)\theta_E + \delta\omega) < \infty$. It follows that $\bar{v} \geq (1 - \delta)\theta_E + \delta\omega > \omega$. $\square$

Proof of Theorem 2

We now complete the proof of Theorem 2.

Proof. Let $\bar{v}, B$ be the highest agent value and the efficient frontier function, as in the statement of Lemma 11. In view of Lemmata 11 and 13, all that remains is to show that the function B is strictly convex on $[\omega, \delta\bar{v} + (\theta_E - \underline{\theta})]$, and that the equilibrium value set is in fact $\hat{\mathcal{E}} := \{(v, b) \in [0, \bar{v}] \times \mathbb{R} : B(v) \leq b \leq m_0 v\} = \{(v, b) \in [0, \bar{v}] \times \mathbb{R} : b \geq B(v), \pi(v, b) \geq 0\}$.

For strict convexity, it suffices to note the following:

- There is some $v_0 \in (\omega, \delta\bar{v} + (1 - \delta)\omega)$ such that B is strictly convex on (ω, v_0) .
- For every $v \in (\omega, \delta\bar{v} + (1 - \delta)\omega)$, there is some $v_P(v) \in (\omega, v)$ such that B is strictly convex on $[\omega, v]$ if it is strictly convex on $[\omega, v_P(v)]$.

Indeed, these together imply that

$$\max \{v \in [\omega, \delta\bar{v} + (1 - \delta)\omega] : B \text{ is strictly convex on } [\omega, v]\},$$

which must exist, is equal to $\delta\bar{v} + (1 - \delta)\omega$. Lemma 12 guarantees the former condition, and the latter condition comes from the functional form of B , letting $v_P(v) := \frac{v - \omega - (1 - \delta)\underline{\theta}}{\delta}$. Strict convexity obtains.

Now, we want to show that $\mathcal{E}^* = \hat{\mathcal{E}}$. To that end, let us show the following claim concerning $B = B_{\mathcal{E}^*}$:

$$\text{ext}\hat{\mathcal{E}} \subseteq \{(\omega, m_0\omega)\} \cup \text{graph}(B).$$

Indeed, consider any $(v, b) \in \text{ext}\hat{\mathcal{E}}$. As such a point lies in the convex hull of $\{(v, B(v)), (v, m_0 v)\}$, extremeness implies that $b \in \{B(v), m_0 v\}$. There is nothing to show if $b = B(v)$, so focus on the case of $b = m_0 v$. Next, extremeness implies that $v \in \{0, \bar{v}\}$, as $(v, m_0 v)$ is a convex combination of $(\bar{v}, m_0 \bar{v})$ and $(0, 0)$. If either $v = 0$ or $\bar{v} > \omega$, then (appealing to Lemma 13) (v, b) belongs to the graph of B . The only remaining possibility is that $b = m_0 v$ and $v = \bar{v} = \omega$. This confirms the claim.

As the graph of B is, by definition, contained in $\mathcal{E}^*$, all that remains is to show that $(\omega, m_0\omega) \in \mathcal{E}^*$. Indeed, the claim would then tell us that

$$\text{ext}\hat{\mathcal{E}} \subseteq \mathcal{E}^* \subseteq \hat{\mathcal{E}},$$

which would prove the result because $\mathcal{E}^*$ and $\hat{\mathcal{E}}$ are both compact and convex.

To see that $(\omega, m_0\omega) \in \mathcal{E}^*$, consider the following strategy profile:

- Whenever the principal delegates to the agent, the agent adopts the project if it is good, and adopts it with probability $\frac{h(\bar{\theta}-c)}{(1-h)(c-\bar{\theta})}$ if it is bad.
- The principal initially delegates. If the principal froze last period, she freezes. If the principal delegated last period, and the agent did not adopt that project, the principal delegates next period. If the principal has always delegated, and the agent adopts a project this period, then the principal freezes next period with probability $\frac{(1-\delta)\bar{\theta}}{\delta\omega}$, and delegates next period with complementary probability.³⁶

One can verify directly that the above is an equilibrium (with the principal always indifferent, and the agent indifferent whenever facing a bad project) which generates a payoff of exactly $(\omega, m_0\omega)$, as required. $\square$

10 APPENDIX: Delayed Differential Equation

Counting time in different units, we may normalize $r = 1$ (or, equivalently, interpret η as $\frac{\eta}{r}$ in all that follows).

Taking a change of variables, from agent value v to account balance $x = \frac{v-\omega}{\bar{\theta}}$, the following system of equations describes the frontier of the equilibrium value set:

$$\begin{aligned} (1 + \eta)b(x) &= \eta b(x - 1) + xb'(x) \text{ for } x > 0, \\ b(x) &= 0 \text{ for } x \leq 0. \end{aligned}$$

Proposition 4. *Consider the above system of equations. For any $\alpha \in \mathbb{R}$, there is a unique solution $b^{(\alpha)}$ to the above system with $b^{(\alpha)}(1) = \alpha$. Moreover $b^{(\alpha)} = \alpha b^{(1)}$.*

Letting $b = b^{(1)}$:

1. $b(x) = x^{1+\eta}$ for $x \in [0, 1]$.
2. b is twice-differentiable on $(0, \infty)$ and globally C^1 on $\mathbb{R}$.
3. b is strictly convex on $(0, \infty)$ and globally convex on $\mathbb{R}$. In particular, b is unbounded.
4. b is strictly increasing and strictly log-concave on $(0, \infty)$.

³⁶That $\frac{(1-\delta)\bar{\theta}}{\delta\omega} \in [0, 1]$ follows from Assumption 1.

Proof. First consider the same equation on a smaller domain,

$$(1 + \eta)b(x) = xb'(x) \text{ for } x \in (0, 1].$$

Taking the standard solution of a first-order linear ODE, the full family of solutions is $\{b^{(\alpha,1)}\}_{\alpha \in \mathbb{R}}$, where $b^{(\alpha,1)}(x) = \alpha x^{1+\eta}$ for $x \in (0, 1]$.

Now, given a particular partial solution $b : (-\infty, z] \rightarrow \mathbb{R}$ up to $z > 0$, there is a unique solution to the first-order linear differential equation $\hat{b} : [z, z+1] \rightarrow \mathbb{R}$ given by

$$\hat{b}'(x) = \frac{1+\eta}{x}\hat{b}(x) - \frac{\eta}{x}b(x-1).$$

Proceeding recursively, there is a unique solution to the given system of equations for each α . Moreover, since multiplying any solution by a constant yields another solution, uniqueness implies $b^{(\alpha)} = \alpha b^{(1)}$. Now let $b := b^{(1)}$.

We have shown that $b(x) = x^{1+\eta}$ for $x \in [0, 1]$, from which it follows readily that b is $C^{1+\lfloor \eta \rfloor}$ on $(-\infty, 1)$,

Given $x > 0$, for small ϵ ,

$$\begin{aligned} (x+\epsilon) \frac{b'(x+\epsilon) - b'(x)}{\epsilon} &= \frac{1}{\epsilon}(x+\epsilon)b'(x+\epsilon) - \frac{1}{\epsilon}xb'(x) - b'(x) \\ &= \frac{1}{\epsilon} \left[(1+\eta)b(x+\epsilon) - \eta b(x+\epsilon-1) \right] - \frac{1}{\epsilon} \left[(1+\eta)b(x) - \eta b(x-1) \right] - b'(x) \\ &= \eta \left[\frac{b(x+\epsilon) - b(x)}{\epsilon} - \frac{b(x-1+\epsilon) - b(x-1)}{\epsilon} \right] + \left[\frac{b(x+\epsilon) - b(x)}{\epsilon} - b'(x) \right] \\ &\xrightarrow{\epsilon \rightarrow 0} \eta[b'(x) - b'(x-1)] + 0. \end{aligned}$$

So b is twice differentiable at $x > 0$ with $b''(x) = \frac{\eta}{x}[b'(x) - b'(x-1)]$.

Let $\bar{x} := \sup\{x > 0 : b'|_{(0,x]} \text{ is strictly increasing}\}$. We know $\bar{x} \geq 1$, from our explicit solution of b up to 1. If $\bar{x}$ is finite, then $b'(\bar{x}) > b'(\bar{x}-1)$. But then $b''(\bar{x}) = \frac{\eta}{\bar{x}}[b'(\bar{x}) - b'(\bar{x}-1)] > 0$, so that b' is strictly increasing in some neighborhood of $\bar{x}$, contradicting the maximality of $\bar{x}$. So $\bar{x} = \infty$, and our convexity result obtains. From that and $b'(0) = 0$, it is immediate that b is strictly increasing on $(0, \infty)$.

Lastly, let $f := \log b|_{(0,\infty)}$. Then $f(x) = (1 + \eta) \log x$ for $x \in (0, 1]$, and for $x \in (1, \infty)$,

$$\begin{aligned}
(1 + \eta)e^{f(x)} &= \eta e^{f(x-1)} + x e^{f(x)} f'(x) \\
\implies (1 + \eta) &= \eta e^{f(x-1)-f(x)} + x f'(x) \\
\implies 0 &= \eta e^{f(x-1)-f(x)} [f'(x-1) - f'(x)] + f'(x) + x f''(x) \\
\implies -x f''(x) &= \eta e^{f(x-1)-f(x)} [f'(x-1) - f'(x)] + f'(x) \\
&\geq \eta e^{f(x-1)-f(x)} [f'(x-1) - f'(x)], \text{ since } f = \log b \text{ is increasing.}
\end{aligned}$$

The same contagion argument will work again. If f has been strictly concave so far, then $f'(x) < f'(x-1)$, in which case $-x f''(x) > 0$ and f will continue to be concave. Since we know $f|_{(0,1]}$ is strictly concave, it follows that f is globally such. $\square$

The first point of the following proposition shows that the economically relevant boundary condition of our DDE uniquely pins down the solution b for any given account cap. The second point shows that as the account cap increases, so does the number of bad projects (in expected discounted terms) anticipated at the cap.

Proposition 5. *For any $\bar{x} > 0$*

1. *There is a unique $\alpha = \alpha(\bar{x}) > 0$ such that $b^{(\alpha)}(\bar{x}) = 1 + b^{(\alpha)}(\bar{x} - 1)$.*
2. *$b^{\alpha(\bar{x})}(\bar{x})$ is increasing in $\bar{x}$.*

Proof. The first part is immediate, with $\alpha = \frac{1}{b^{(1)}(\bar{x}) - b^{(1)}(\bar{x}-1)}$.

For the second part, notice that $\frac{b(\bar{x})}{b(\bar{x}-1)}$ is decreasing in $\bar{x}$ because b is log-concave. Then,

$$b^{\alpha(\bar{x})}(\bar{x}) = \alpha(\bar{x}) b(\bar{x}) = \frac{b(\bar{x})}{b(\bar{x}) - b(\bar{x}-1)} = \frac{1}{1 - \frac{b(\bar{x}-1)}{b(\bar{x})}}$$

is increasing in $\bar{x}$. $\square$

11 APPENDIX: Comparative Statics

In this section, we prove Proposition 3.

For any parameters $\eta, \bar{\theta}, \underline{\theta}, c$ satisfying Assumption 1, and for any balance and bad projects x, b

satisfying $x \geq b > 0$, define the associated profit

$$\begin{aligned}
\hat{\pi}(x, b | \eta, \bar{\theta}, \underline{\theta}, c) &:= \eta(\bar{\theta} - \underline{\theta}) + \underline{\theta}x - c \left[\frac{\eta(\bar{\theta} - \underline{\theta}) + \underline{\theta}x - \underline{\theta}b}{\bar{\theta}} + b \right] \\
&= \left(1 - \frac{c}{\bar{\theta}}\right) [\eta(\bar{\theta} - \underline{\theta}) + \underline{\theta}x] - c \left(1 - \frac{\underline{\theta}}{\bar{\theta}}\right) b \\
&= (\bar{\theta} - c)\eta + \left(1 - \frac{c}{\bar{\theta}}\right) \underline{\theta}(x - \eta) - c \left(1 - \frac{\underline{\theta}}{\bar{\theta}}\right) b.
\end{aligned}$$

For reference, we compute the following derivatives of profit:

$$\begin{aligned}
\frac{\partial \hat{\pi}}{\partial \bar{\theta}} &= \eta - c \underline{\theta} \frac{-1}{\bar{\theta}^2} [x - b - \eta] = \left(1 - \frac{c \underline{\theta}}{\bar{\theta}^2}\right) \eta + \frac{c \underline{\theta}}{\bar{\theta}^2} (x - b) > 0. \\
\frac{\partial \hat{\pi}}{\partial \underline{\theta}} &= \left(1 - \frac{c}{\bar{\theta}}\right) (x - \eta) + c \frac{1}{\bar{\theta}} b, \text{ which implies} \\
(\bar{\theta} - \underline{\theta}) \frac{\partial \hat{\pi}}{\partial \underline{\theta}} + \hat{\pi} &= (\bar{\theta} - c)\eta + \left(1 - \frac{c}{\bar{\theta}}\right) \bar{\theta}(x - \eta) - 0b = (\bar{\theta} - c)x > 0. \\
\frac{\partial \hat{\pi}}{\partial c} &= -\frac{1}{\bar{\theta}} [\eta(\bar{\theta} - \underline{\theta}) + \underline{\theta}x + (\bar{\theta} - \underline{\theta})b] < 0.
\end{aligned}$$

Fix parameters $(\bar{\theta}, \underline{\theta}, c)$, and let $\bar{x}^*$ be as delivered in Theorem 3. We first show that slightly raising either of $\bar{\theta}, \underline{\theta}$ or slightly lowering c weakly raises the cap, strictly if $\bar{x}^* > 0$.

- If $\bar{x}^* = 0$, there is nothing to show, so assume $\bar{x}^* > 0$ henceforth. Notice that the expected discounted number of bad projects when at the cap depends only on η and the size of the cap. By Theorem 3, $\hat{\pi}(\bar{x}^*, b | \eta, \bar{\theta}, \underline{\theta}, c) = 0$.
- Consider slightly raising $\bar{\theta}$ or $\underline{\theta}$, or slightly lowering c . By the above derivative computations, the profit of the DCB contract with cap $\bar{x}^*$ is strictly positive.
- For any of the above considered changes, the DCB contract with cap $\bar{x}^*$ has strictly positive profits. Appealing to continuity, a slightly higher cap still yields positive profits under the new parameters, and is therefore consistent with equilibrium by Proposition 2. Then, appealing to Theorem 3 again, the cap associated with the new parameters is strictly higher than $\bar{x}^*$.

Now we consider comparative statics in the profit-maximizing initial account balance. We will fix parameters $(\bar{\theta}, \underline{\theta}, c)$, and consider raising either of $\bar{\theta}, \underline{\theta}$ or lowering c .

- Again, if $\bar{x}^* = 0$ at the original parameters, there's nothing to check, so assume $\bar{x}^* > 0$.
- Let $\check{b}$ be some solution to the DDE in Section 10, so that the expected discounted number of bad projects at a given cap $\bar{x}$ and balance x is $b(x | \bar{x}) = \frac{\check{b}(x)}{\check{b}(\bar{x}) - \check{b}(\bar{x} - 1)}$. Because $\check{b}$ is strictly

increasing and strictly convex (by work in Section 10), we know that $\frac{\partial}{\partial x}b(x|\bar{x})$ is strictly decreasing in $\bar{x}$. Therefore, by our comparative statics result for the cap, the parameter change results in a global strict decrease of $\frac{\partial}{\partial x}b(x|\bar{x}^*)$.

- By the form of $\hat{\pi}$ and by convexity of $b(\cdot|\bar{x})$, the unique optimal initial balance is the balance at which $\frac{\partial}{\partial x}b(x|\bar{x})$ is equal to

$$m_0 = \frac{\underline{\theta}\left(1 - \frac{c}{\underline{\theta}}\right)}{c\left(1 - \frac{\underline{\theta}}{\bar{\theta}}\right)} = \frac{\underline{\theta}(\bar{\theta} - c)}{c(\bar{\theta} - \underline{\theta})},$$

which increases with the parameter change.

- As m_0 increases and $\frac{\partial}{\partial x}b(x|\bar{x}^*)$ decreases (at each x) with the parameter change, the optimal balance x^* must increase (given convexity) to satisfy the first-order condition $\frac{\partial}{\partial x}b(x|\bar{x})\Big|_{x=x^*} = m_0$.

References

- Abdulkadiroglu, Atila and Kyle Bagwell. 2013. "Trust, Reciprocity, and Favors in Cooperative Relationships." *American Economic Journal: Microeconomics* 5 (2):213–59.
- Abreu, D., D. Pearce, and E. Stacchetti. 1990. "Toward a Theory of Discounted Repeated Games with Imperfect Monitoring." *Econometrica* 58 (5):1041–63.
- Alonso, R. and N. Matouschek. 2007. "Relational delegation." *RAND Journal of Economics* 38 (4):1070–1089.
- Ambrus, Attila and Georgy Egorov. 2017. "Delegation and nonmonetary incentives." *Journal of Economic Theory* .
- Armstrong, M. and J. Vickers. 2010. "A Model of Delegated Project Choice." *Econometrica* 78 (1):213–44.
- Baumol, W. 1968. "On the Theory of Oligopoly." *Economica* 25 (99):187–98.
- Biais, B., T. Mariotti, J. Rochet, and S. Villeneuve. 2010. "Large Risks, Limited Liability, and Dynamic Moral Hazard." *Econometrica* 78 (1):73–118.
- Bird, D. and A. Frug. 2017. "Dynamic Nonmonetary Incentives." Working paper, Tel Aviv University & Pompeu Fabra.
- Casella, A. 2005. "Storable Votes." *Games and Economic Behavior* 51:391–419.
- Clementi, G. and H. Hopenhayn. 2006. "A Theory of Financing Constraints and Firm Dynamics." *Quarterly Journal of Economics* 121 (1):1131–1175.
- DeMarzo, Peter M and Michael J Fishman. 2007. "Optimal long-term financial contracting." *Review of Financial Studies* 20 (6):2079–2128.
- Edmans, Alex, Xavier Gabaix, Tomasz Sadzik, and Yuliy Sannikov. 2012. "Dynamic CEO compensation." *Journal of Finance* 67 (5):1603–1647.
- Espino, E., J. Kozlowski, and J. Sanchez. 2016. "Investment and Incentives in Partnerships." Working paper, Torcuato Di Tella, NYU, & FRB of St. Louis.
- Fong, Yuk-fai and Jin Li. 2017. "Relational contracts, limited liability, and employment dynamics." *Journal of Economic Theory* 169:270–293.
- Frankel, A. 2014. "Aligned Delegation." *American Economic Review* 104 (1):66–83.
- Frankel, Alex. 2016a. "Delegating multiple decisions." *American Economic Journal: Microeconomics* 8 (4):16–53.

- Frankel, Alexander. 2016b. “Discounted quotas.” *Journal of Economic Theory* 166:396–444.
- Fudenberg, Drew, David Levine, and Eric Maskin. 1994. “The folk theorem with imperfect public information.” *Econometrica* :997–1039.
- Guo, Y. and J. Hörner. 2015. “Dynamic Mechanisms without Money.” Working paper, Northwestern University & Yale University.
- Guo, Yingni. 2016. “Dynamic delegation of experimentation.” *American Economic Review* 106 (8):1969–2008.
- Harris, M. and B. Holmström. 1982. “A theory of wage dynamics.” *Review of Economic Studies* 49 (3):315–333.
- Hauser, C. and H. Hoppenhayn. 2008. “Trading Favors: Optimal Exchange and Forgiveness.” Working paper, Collegio Carlo Alberto & UCLA.
- Holmström, B. 1984. “On the Theory of Delegation.” In *Bayesian Models in Economic Theory*. New York: North-Holland, 115–41.
- Jackson, M. and H. Sonnenschein. 2007. “Overcoming Incentive Constraints by Linking Decisions.” *Econometrica* 75 (1):241–57.
- Lazear, E. P. 1981. “Agency, earnings profiles, productivity, and hours restrictions.” *American Economic Review* 71 (4):606–620.
- Li, Jin, Niko Matouschek, and Michael Powell. 2017. “Power dynamics in organizations.” *American Economic Journal: Microeconomics* 9 (1):217–241.
- Mailath, G. and L. Samuelson. 2006. *Repeated Games and Reputations: Long-Run Relationships*. New York: Oxford University Press.
- Malenko, A. 2016. “Optimal Dynamic Capital Budgeting.” Working paper, MIT.
- Möbius, M. 2001. “Trading Favors.” Working paper, Harvard University.
- Nayyar, Shivani. 2009. *Essays on repeated games*. Princeton University.
- Ray, Debraj. 2002. “The Time Structure of Self-Enforcing Agreements.” *Econometrica* 70 (2):547–582.
- Rubinstein, Ariel. 1979. “An optimal conviction policy for offenses that may have been committed by accident.” In *Applied game theory*. Springer, 406–413.
- Spear, S. and S. Srivastava. 1987. “On Repeated Moral Hazard with Discounting.” *Review of Economic Studies* 54 (4):599–617.

Thomas, Jonathan and Tim Worrall. 1990. "Income fluctuation and asymmetric information: An example of a repeated principal-agent problem." *Journal of Economic Theory* 51 (2):367 – 390.