

Reward and Punishment in a Regime Change Game¹

Stephen Morris² Mehdi Shadmehr³

May, 2017

¹We are grateful for comments of Nemanja Antic, Dan Bernhardt, Raphael Boleslavsky, Georgy Egorov, Alex Jakobsen, Dimitri Migrow, Konstantin Sonin, seminar participants at the 2016 Columbia conference on Political Economy, 2017 North American Summer Meeting of the Econometric Society, University of Miami and Queen's University, and the discussion of Alex Debs at the Columbia conference. We have received excellent research assistance from Elia Sartori.

²Department of Economics, Princeton University, Julis Romo Rabinowitz Building, Princeton, NJ 08544. Email: smorris@princeton.edu.

³Department of Economics, University of Calgary, 2500 University Dr., Calgary, AB, T2N 1N4. E-mail: mehdi.shadmehr@ucalgary.ca.

Abstract

We characterize how a leader should allocate rewards to different levels of citizen participation to maximize the likelihood of regime change. Regime change occurs when the aggregate revolutionary effort from all citizens exceeds the uncertain regime's strength, about which citizens have private information. Citizens face a coordination problem in which each citizen has a private (endogenous) belief about the likelihood of regime change. Because optimistic citizens are easier to motivate, the leader faces a screening problem: assigning higher rewards to a level of participation can induce more activity in some citizens, but less in others. We analyze this continuous action global game with screening that arises from the strategic interactions between a leader and citizens. An optimal reward scheme gives rise to a distinct group akin to revolutionary vanguards: a group of more optimistic citizens who all engage in the (endogenous) maximum level of revolutionary activity. Activity level among less optimistic citizens falls smoothly with their level of pessimism.

1 Introduction

Consider the problem of a revolutionary leader who wishes to induce citizen participation in a revolutionary movement. The revolutionary leader must decide what level of participation to solicit, e.g., peaceful demonstration or engaging in violence. In particular, the leader must decide how to allocate psychological (or other) rewards to different levels of participation. The leader's optimal choice depends on the punishments chosen by the existing regime. This paper identifies a leader's optimal reward scheme in this context. We discuss in turn our modeling of this problem, the optimal reward scheme and how our modeling choices and the resulting optimal reward scheme capture evidence on leadership, citizen motivation and participation levels in revolutionary movements.

We analyze a coordination game model of regime change, where there are a continuum of ex-ante identical players (*citizens*) and a continuum of actions (*effort*, capturing the level of participation). The revolution succeeds (*regime change*) if the aggregate effort, summing across citizens, exceeds a critical level, with the critical level depending on the state of the world. Any given effort level gives rise to a punishment (the *cost*). The leader must choose a reward (the *benefit*) for each effort level. The benefits are enjoyed only if there is regime change. There is an upper bound on the level of benefits that the leader can induce for any level of effort. We study this problem in a *global game* model (Carlsson and van Damme 1993) of regime change (Morris and Shin 1998), where the state is observed with a small amount of noise, so that perfect coordination is not possible.

Two forces drive our results. First, there is *strategic uncertainty*. At states where there is regime change, citizens observe heterogeneous signals giving rise to endogenous heterogeneity in the population's beliefs about the likelihood of regime change. Second, there is *screening*. The leader cannot distinguish more pessimistic from more optimistic citizens, and more optimistic citizens are easier to motivate. In particular, if the leader increases the benefit for a given level of effort, it induces higher effort from those who would otherwise have chosen lower levels of effort. But it also induces lower effort from those who would otherwise have chosen higher levels of effort, since they can get the reward at lower cost. The optimal reward scheme trades off these countervailing effects.

The optimal reward scheme has two regions. First, there is an endogenous maximum effort level such that citizens receive the maximum benefit only if they choose that maximum effort (or higher). Second, below the maximum effort level, benefits depend continuously on the level of effort, converging to zero as effort goes to zero and converging to the maximum benefit as effort goes to the maximum level. A mass of more optimistic citizens pool on the maximum effort. A mass of more pessimistic citizens choose zero effort. For intermediate levels of optimism, citizens separate, choosing intermediate effort levels. The result reflects the fact that inducing more effort from the most optimistic citizens (who already exert more effort than others) is very costly in terms of the amount of effort that can be induced from less optimistic citizens; and inducing effort from the most pessimistic citizens generates second order gains for the leader at the expense of first order costs.

There are two key modeling choices. What induces citizens to participate in revolutionary activity? Wood (2003) argues that participation is not based on citizens having

the unreasonable belief that their probability of being pivotal compensates for the cost of participation. Nor is citizens' participation expressive—a psychic benefit that they receive independent of the likelihood of success. Rather, participation is driven by *pleasure in agency*: a psychological benefit that depends on the likelihood of success, but does not depend on the ability to influence that likelihood. Benefits in our model correspond exactly to such pleasure in agency payoffs. But pleasure in agency is not exogenous. If pleasure in agency drives willingness to participate, a key role of leadership is to inspire pleasure in agency. Our model of leadership assumes that a leader influences how much participation is required for a given pleasure in agency, subject to a constraint on the maximum possible pleasure in agency. Through a variety of channels such as speeches, writings, and meetings, the leader creates or manipulates different amounts of pleasure in agency for different levels of anti-regime activities. Our key finding is that there are three distinct groups: a set of *professional* revolutionaries engaging in the (endogenous) maximum level of effort and receiving the maximum level of pleasure in agency; a set of part-time revolutionaries who participate at varying levels with corresponding variation in pleasure in agency; and a set of non-participants. The first group corresponds to the vanguard in Lenin's treatise, *What Is to be Done?* Section 4 provides a detailed discussion of pleasure in agency, participation levels, leadership and revolutionary organization.

We now describe our results in more detail, before discussing additional results that we could attain with our approach, highlighting novel theoretical contributions and discussing the relation to the literature.

We start by analyzing the regime change model with *exogenous* benefits as well as exogenous costs. The state, or strength of the regime, is the (minimum) total amount of citizens' efforts that gives rise to regime change. We show that if the cost function is increasing, optimal effort is increasing in the probability p that a citizen assigns to regime change. This is true whatever the shape of the benefit function. Thus, an optimal effort function will have a minimum effort level e_{min} (when $p = 0$) and a maximum effort level e_{max} (when $p = 1$). When there is common knowledge of the state, there is an equilibrium where all citizens exert effort e_{max} as long as the state is less than e_{max} ; and there is an equilibrium where all citizens exert effort level e_{min} as long as the state is more than e_{min} . Thus, both equilibria exist at states between e_{min} and e_{max} .

Now suppose that there is incomplete information about the state. We focus on the case where the state is uniformly distributed and citizens observe noisy signals of the state (as is well known in the global game literature, the results extend to any smooth prior as long as citizens' signals are sufficiently accurate). We exploit a key statistical property that we call *uniform threshold belief*. For a given threshold state, a citizen's threshold belief is the probability that he assigns to the true state being below that threshold. At any given threshold state, there will be a realized distribution of threshold beliefs in the population. Guimaraes and Morris (2007) showed that this threshold belief distribution is uniform. An intuition for this surprising result is the following. At any given threshold state and for any given citizen, we can identify the citizen's rank in the threshold belief distribution, i.e., the proportion of citizens with higher signals. This rank is necessarily uniformly distributed: Since a citizen does not know his rank, his belief about whether the true state is below the

threshold state is uniform.

We consider monotone strategy profiles of the incomplete information game, where each citizen's effort is (weakly) decreasing in his signal. Any such strategy profile gives rise to a unique *regime change threshold*, such that there will be regime change if the true state is below that threshold. Using the uniform threshold belief property, we show that in a monotone equilibrium, the total effort at the regime change threshold is the integral of the optimal effort function between $p = 0$ and $p = 1$ (the *uniform threshold*). This implies that there is a unique monotone equilibrium where the regime changes only when the state is above this uniform threshold.

Having established results for the case of exogenous benefit and cost functions, we then consider the problem of the leader choosing the benefit function, subject to the constraint that there is a maximum possible level of benefit. What benefit function would he choose and what would be the implications for the level and distribution of effort choices in the population? Using the equilibrium characterization described above, we show that this problem reduces to finding the benefit function that maximizes the uniform threshold. But this in turn reduces to a screening problem, where the leader chooses effort levels to maximize the threshold, subject to the constraint that the effort levels can be induced in an incentive compatible way using the benefit function. Our main result corresponds to the solution of that screening problem.

This paper makes multiple methodological contributions, each of which is of independent interest, and could be used to study further questions of methodological and substantive interest. One contribution is to solve a coordination game of regime change with continuous actions and exogenous payoffs and characterize the essentially unique monotone equilibrium. The equilibrium characterization can be used in a wide range of economic applications that have been modelled as global games with regime change (almost always with binary actions). We could study a variety of other comparative statics questions using this characterization. What would happen if we endogenized the costs chosen by the existing regime (instead of the benefits chosen by the revolutionary leader)? What if both were chosen endogenously? How would our analysis with exogenous costs and benefits change if we replaced our continuous action assumption with binary actions, the case usually studied in the global games literature and a natural assumption in some contexts? How is the optimal reward scheme changed if the leader is restricted to binary participation levels? A second contribution is to analyze a setting where coordination and screening are intertwined; in other words, a screening problem arises only because of the citizen heterogeneity introduced by strategic uncertainty, and despite the fact that citizens are ex ante identical. What would happen in other settings where this interaction naturally occurred? Finally, as an ingredient to our solution, we solve a particular screening problem. What happens in an environment without transfers where a principal can choose a reward scheme with no cost and only constrained by the requirement that rewards belong to a closed interval? Such screening problems will naturally arise with exogenous agent heterogeneity independent of any strategic component. Among many contexts where this principal agent problem might arise is one where the principal is a fundraiser choosing recognition for donors of unknown types making heterogeneous contributions. All these alternative directions are

left for future work.

Our paper builds on a literature on global games in general, and applications in political economy in particular (Bueno de Mesquita 2010; Edmond 2013; Chen and Suen 2016; Quigley and Toscani 2016; Tyson and Smith 2016; Egorov and Sonin 2017; Shadmehr 2017). A distinctive aspect of this paper is that players can choose continuous actions. This case has been studied in an abstract setting by Frankel, Morris and Pauzner (2003). Guimaraes and Morris (2007) identify the uniform threshold belief property and use it to provide a tractable solution (their application is to currency attacks). Our work builds heavily on this insight. They study a supermodular payoff setting where there is a unique equilibrium that is also dominance solvable. Our problem does not have supermodular payoffs, but it gives rise to monotonic strategies in equilibrium and we end up with their characterization, although we can establish only a weaker uniqueness result (among all monotone equilibria).¹ Our analysis of the screening problem exploits classic arguments from the screening literature. Guesnerie and Laffont (1984) is a key early reference on a rich class of screening problems that embeds monopoly problems of choosing quality (Mussa and Rosen 1978) or quantity (Maskin and Riley 1984) and the government’s regulation of a monopolist (Baron and Myerson 1982). In our problem, a principal gives “benefits” to an agent in exchange for “effort”. Our screening problem is non-standard because agent utility is not linear in effort and the principal’s “budget” of rewards is not “smooth”, i.e., that rewards up to a level are free to the principal, but higher rewards have infinite costs.

Mookherjee and Png (1994) study the problem of choosing the optimal likelihood of detection and optimal punishment schedule by an authority. Their key result is that the marginal punishment should remain lower than marginal harms of crime for low crime levels; In fact, if monitoring is sufficiently costly, then a range of less harmful acts should be legalized—bunching at the bottom. We solve a problem where strategic global games analysis is combined with screening, and focus on the design of benefits. Absent coordination problem among citizens (agents), their analysis is related to the design of costs in our setting. However, when designing costs, our regime (principal) would aim to minimize the aggregate actions, while their authority aims to maximize social welfare consisting of the criminal’s benefits from crime minus the crime’s exogenous harm and the costs of law enforcement. Recent work of Shen and Zou (2016) also addresses an interaction between strategic behavior and screening. They consider a binary action coordination game but allow players to receive a stimulus by self-selection. This contingent policy intervention allows coordination on a desirable outcome at minimal cost. Laffont and Robert’s (1996) and Hartline’s (2016, Ch. 8) analysis of optimal auction with financial constraints features bunching at the top. In their settings, because no buyer can bid above a known budget (e.g., due to an exogenous liquidity constraint), the seller cannot separate eager buyers.

We describe and solve the exogenous benefits case in Section 2, describe and solve the optimal reward problem in Section 3, and motivate our model and results in Section 4.

¹Assuming that an optimal benefit scheme is increasing in effort is sufficient for supermodularity and hence uniqueness among all equilibria. However, when designing the benefits, benefit schemes are endogenous, and such an assumption is ex-ante restrictive.

2 Exogenous Rewards and Punishments

2.1 Model

There is a continuum of citizens, indexed by $i \in [0, 1]$, who must decide how much effort $e \geq 0$ to contribute to regime change. Exerting effort e costs $C(e)$, independent of whether there is regime change, and gives a benefit $B(e)$ only if there is regime change. Thus, if a citizen thinks that the regime change occurs with probability p , his expected payoff from choosing effort e is

$$pB(e) - C(e).$$

Regime change occurs if and only if the total contribution $\int e_i di$ is greater than or equal to the regime's strength $\theta \in \mathbb{R}$. Thus, if citizens choose effort levels $(e_j)_{j \in [0,1]}$, the payoff to a citizen from effort level e_i is

$$B(e_i) I_{\int e_j dj \geq \theta} - C(e_i).$$

The regime's strength θ is uncertain, and citizens have an improper common prior that θ is distributed uniformly on $\mathbb{R}$. In addition, each citizen $i \in [0, 1]$ receives a noisy private signal $x_i = \theta + \nu_i$ about θ , where θ and the ν_i 's are independently distributed with $\nu_i \sim f(\cdot)$.²

2.2 A Key Statistical Property: Uniform Threshold Belief

Fix any state $\hat{\theta}$. A citizen's *threshold belief* about $\hat{\theta}$ is the probability that he assigns to the event that $\theta \leq \hat{\theta}$. The *threshold belief distribution* at $\hat{\theta}$ is the distribution of threshold beliefs about $\hat{\theta}$ in the population when the true state is $\hat{\theta}$. We will show that the threshold belief distribution is always uniform, regardless of the state $\hat{\theta}$ and the distribution of noise.

To state the uniform threshold belief property formally, write $H(\cdot|\hat{\theta})$ for the cdf of the threshold belief distribution at $\hat{\theta}$. Thus, the proportion of citizens with threshold belief about $\hat{\theta}$ below p when the true state is $\hat{\theta}$ is $H(p|\hat{\theta})$. The key to establishing the result is the observation that, because θ is distributed uniformly so that there is no prior information about θ , one can consider θ as a signal of x_i , writing $\theta = x_i - \nu_i$. So the threshold belief about $\hat{\theta}$ for a citizen observing x_i is simply the probability that $x_i - \nu_i \leq \hat{\theta}$, or $1 - F(x_i - \hat{\theta})$. Thus, a citizen observing x_i has a threshold belief less than or equal to p if $1 - F(x_i - \hat{\theta}) \leq p$, or $x_i \geq \hat{\theta} + F^{-1}(1 - p)$. But if the true state is $\hat{\theta}$, the proportion of citizens observing a signal greater than or equal to $\hat{x}$ will be the proportion of citizens with $\hat{\theta} + \nu_i \geq \hat{x}$, or $1 - F(\hat{x} - \hat{\theta})$. Combining these observations, we have that $H(p|\hat{\theta}) =$

²The assumption of an improper common prior is standard in this literature. At the cost of further notation, one can replace this assumption with a prior that is uniformly distributed on a sufficiently large bounded interval. More importantly, results proved under the uniform prior assumption can also be reproduced with an arbitrary smooth prior when the noise is sufficiently small: see Carlsson and van Damme (1993), Morris and Shin (2003) and Frankel, Morris and Pauzner (2003). All our results thus hold with general priors in the limit as noise goes to zero.

$1 - F(\hat{x} - \hat{\theta})$ where $\hat{x} = \hat{\theta} + F^{-1}(1 - p)$. Substituting for $\hat{x}$, we have

$$H(p|\hat{\theta}) = 1 - F(\hat{x} - \hat{\theta}) = 1 - F(F^{-1}(1 - p)) = 1 - (1 - p) = p.$$

We have shown the uniform threshold belief property:

Lemma 1. (*Guimaraes and Morris (2007)*) *The threshold beliefs distribution is always uniform on $[0, 1]$, so that $H(p|\hat{\theta}) = p$ for all $p \in [0, 1]$ and $\hat{\theta} \in \mathbb{R}$.*

To gain intuition for the result, suppose that noise itself was uniformly distributed on the interval $[0, 1]$. Now if the true state was $\hat{\theta}$, a citizen observing x_i in the interval $[\hat{\theta}, \hat{\theta} + 1]$ would have threshold belief $\hat{\theta} + 1 - x_i$ and we would have uniform threshold belief. But now if the noise had some arbitrary distribution, we could do a change of variable replacing the level of a citizen's signal with its percentile in the distribution. Because citizens do not know their own percentiles, the same argument then goes through.

2.3 Equilibrium

Our results will depend crucially on how citizens would behave if they assigned a fixed probability to the success of the revolution. It is convenient to state our results first depending only on the monotonicity and continuity of the solution to this problem, and only then identify conditions on the exogenous cost function $C(\cdot)$ and the exogenous (in this section) or endogenous (in the next section) benefit function $B(\cdot)$ that imply monotonicity and continuity.

Definition 1. The *optimal effort correspondence* is

$$e^*(p) = \arg \max_{e \geq 0} pB(e) - C(e);$$

we write

$$e_{min} = \min(e^*(0)) \text{ and } e_{max} = \max(e^*(1)).$$

We maintain the assumption that the maximum and minimum exist (and are finite).

Definition 2. The optimal effort correspondence is *weakly increasing* if $p_2 > p_1$ and $e_i \in e^*(p_i)$, $i \in \{1, 2\}$, imply $e_2 \geq e_1$.

That is, a citizen who is strictly more optimistic about the likelihood of success than another will put in at least as much effort as him. Note that if e^* is weakly increasing, it is almost everywhere single-valued. We first consider the complete information case where there is common knowledge of θ .

Proposition 1. *Suppose that the optimal effort correspondence is weakly increasing and that there is complete information. There is an equilibrium with regime change if $\theta \leq e_{max}$, e.g., with all players choosing effort level e_{max} ; and there is an equilibrium without regime change if $\theta > e_{min}$, e.g., with all players choosing effort level e_{min} . Thus, there are three cases: if $\theta \leq e_{min}$, there are only equilibria with regime change; if $\theta > e_{max}$, there are only equilibria without regime change; and if $e_{min} < \theta \leq e_{max}$, then there are equilibria with regime change and equilibria without regime change.*

Proof. If there is regime change in equilibrium, all citizens must choose $e \in e^*(1)$; this can give rise to regime change only if $e_{max} \geq \theta$. If there is no regime change in equilibrium, all citizens must choose $e \in e^*(0)$; this can give rise to no regime change only if $e_{min} < \theta$. $\square$

We now consider the incomplete information game where a citizen's only information is the signal that he receives, maintaining the assumption that the optimal effort correspondence is weakly increasing. Recall that because $e^*(p)$ is a weakly increasing correspondence, it is single valued almost everywhere. A strategy for a citizen i is a mapping $s_i : \mathbb{R} \rightarrow \mathbb{R}_+$, where $s_i(x_i)$ is the effort level of citizen i when he observes signal x_i . Each strategy profile $(s_i)_{i \in [0,1]}$ will give rise to aggregate behavior

$$\hat{s}(\theta) = \int_{i=0}^1 \left(\int_{\nu_i=-\infty}^{\infty} s_i(\theta + \nu_i) f(\nu_i) d\nu_i \right) di.$$

If all citizens follow weakly decreasing strategies, then $\hat{s}$ is weakly decreasing, and there is a unique threshold θ^* such that $\hat{s}(\theta^*) = \theta^*$. Then, $\hat{s}(\theta) \geq \theta$ and there will be a regime change if and only if $\theta \leq \theta^*$. Thus, a citizen observing a signal x_i will assign probability $G(\theta^*|x_i)$ to the event that $\theta \leq \theta^*$ and thus to regime change. We conclude that each citizen must be following the strategy with

$$s^*(x_i) = e^*(G(\theta^*|x_i)).$$

Letting $p = G(\theta^*|x_i)$, and recalling that $H(p|\theta^*)$ is the cdf of $G(\theta^*|x_i)$ conditional on θ^* , this implies that the aggregate effort of citizens when the true state is θ^* is

$$\hat{s}(\theta^*) = \int_{p=0}^1 e^*(p) dH(p|\theta^*).$$

By Lemma 1, we know that $H(p|\theta^*)$ is a uniform distribution. We also know that $\theta^* = \hat{s}(\theta^*)$. Thus,

$$\theta^* = \hat{s}(\theta^*) = \int_{p=0}^1 e^*(p) dp.$$

Summarizing this analysis, we have:

Lemma 2. *If the optimal effort correspondence e^* is weakly increasing, then there is a unique monotonic equilibrium where all citizens follow the essentially unique common strategy*

$$s^*(x_i) = e^* (1 - F(x_i - \theta_{e^*})). \quad (1)$$

In this equilibrium, there is regime change whenever $\theta \leq \theta_{e^}$, where the regime change threshold θ_{e^*} is given by*

$$\theta_{e^*} = \int_{p=0}^1 e^*(p) dp. \quad (2)$$

Guimaraes and Morris (2007) proved the stronger result that the corresponding equilibrium was unique among all equilibria, but under assumptions implying supermodular payoffs. The above proposition holds under weaker assumptions, requiring only that $e^*(p)$ is increasing in p . We have proved uniqueness within the class of monotonic equilibria, but leave open the question of whether non-monotonic equilibria exist.

We can provide a clean characterization of the regime change threshold if e^* is single-valued and continuous. Note that, under these assumptions, $e^*(0) = \{e_{min}\}$ and $e^*(1) = \{e_{max}\}$.

Lemma 3. *If the optimal effort correspondence e^* is single-valued, continuous and weakly increasing, and $e_{min} = 0$, then the regime threshold in the unique monotonic equilibrium is given by*

$$\theta_{e^*} = \int_{e=0}^{e_{max}} \left(1 - \frac{C'(e)}{B'(e)}\right) de. \quad (3)$$

Proof. By assumption, $e^*(p)$ is a weakly increasing function. Whenever $e^*(p)$ is strictly increasing, we have $pB'(e^*(p)) = C'(e^*(p))$. Let $[p_1, p_2]$ be an interval on which e^* is strictly increasing—and continuous. In this case,

$$\begin{aligned} \int_{p=p_1}^{p_2} e^*(p) dp &= [p e^*(p)]_{p=p_1}^{p_2} - \int_{p=p_1}^{p_2} p [e^*]'(p) dp \\ &= p_2 e^*(p_2) - p_1 e^*(p_1) - \int_{e=e^*(p_1)}^{e^*(p_2)} \frac{C'(e)}{B'(e)} de, \end{aligned}$$

where the last inequality uses the change of variables $e = e^*(p)$, and the derivative of $B'(e)$ at $e^*(p_2)$ is the left derivative and at $e^*(p_1)$ is the right derivative.

Now let $[p_2, p_3]$ be an interval on which e^* is constant. In this case,

$$\begin{aligned} \int_{p=p_2}^{p_3} e^*(p) dp &= (p_3 - p_2) e^*(p_3) \\ &= p_3 e^*(p_3) - p_2 e^*(p_2). \end{aligned}$$

We conclude that

$$\int_{p=p_1}^{p_3} e^*(p) dp = p_3 e^*(p_3) - p_1 e^*(p_1) - \int_{e=e^*(p_1)}^{e^*(p_2)} \frac{C'(e)}{B'(e)} de.$$

Now, consider a partition of $[0, 1]$, and suppose that $e^*(p)$ is strictly increasing on the intervals $[p_1, p_2], [p_3, p_4], \dots, [p_{2n-1}, p_{2n}]$, where $p_1 < p_2 < \dots < p_{2n}$. Then,

$$\begin{aligned} \int_{p=0}^1 e^*(p) dp &= e_{max} - \sum_{m=1}^n \left(\int_{e=e^*(p_{2m-1})}^{e^*(p_{2m})} \frac{C'(e)}{B'(e)} de \right) \\ &= e_{max} - \int_{e=0}^{e_{max}} \frac{C'(e)}{B'(e)} de \\ &= \int_{e=0}^{e_{max}} \left(1 - \frac{C'(e)}{B'(e)} \right) de. \end{aligned}$$

□

We will use this characterization of the regime change threshold in our main result.

2.4 Optimal Effort

Having shown the existence and characterizations of a unique monotonic equilibrium for a given monotonic optimal effort correspondence, we now identify sufficient conditions on the exogenous cost function $C(\cdot)$ for the existence of such a monotonic optimal effort correspondence. We first present a weak sufficient condition for monotonicity, and then report standard sufficient conditions for continuity as well as monotonicity.

First suppose that punishments are increasing in effort, so that $C(e)$ is increasing.

Proposition 2. *If $C(e)$ is strictly increasing in e , then any selection from optimal effort correspondence is weakly increasing. Thus the unique monotone equilibrium has regime change threshold given by equation (2) of Lemma 2.*

Proof. Let $p_2 > p_1 > 0$, $e_i \in e^*(p_i) = \arg \max_{e \geq 0} p_i B(e) - C(e)$, $i \in \{1, 2\}$. We establish that $e_2 \geq e_1$, by way of contradiction. Suppose not, so that $e_2 < e_1$. First, from the optimality of e_1 and e_2 , we have:

$$p_2 B(e_2) - C(e_2) \geq p_2 B(e_1) - C(e_1) \quad \Leftrightarrow \quad p_2 [B(e_2) - B(e_1)] \geq C(e_2) - C(e_1) \quad (4)$$

$$p_1 B(e_1) - C(e_1) \geq p_1 B(e_2) - C(e_2) \quad \Leftrightarrow \quad C(e_2) - C(e_1) \geq p_1 [B(e_2) - B(e_1)] \quad (5)$$

Because $C(e)$ is strictly increasing, $e_2 < e_1$ implies $C(e_2) < C(e_1)$. Thus, from (5):

$$0 > C(e_2) - C(e_1) \geq p_1 [B(e_2) - B(e_1)] \Rightarrow B(e_2) - B(e_1) < 0,$$

and hence:

$$p_1 [B(e_2) - B(e_1)] > p_2 [B(e_2) - B(e_1)]. \quad (6)$$

However, combining (4) and (5), we have:

$$p_1 [B(e_2) - B(e_1)] \leq C(e_2) - C(e_1) \leq p_2 [B(e_2) - B(e_1)]. \quad (7)$$

Hence, (6) and (7) contradict each other, and hence our assumption that $e_2 < e_1$ must be false. The proposition is then implied by Lemma 2. $\square$

Proposition 2 will be used to establish the existence and characterization of a monotone equilibrium for any reward scheme B .

We can also give completely standard conditions for continuity:

Proposition 3. *Suppose that costs and benefits are (1) twice continuously differentiable, (2) zero with zero effort ($C(0) = B(0) = 0$), (3) strictly increasing ($C'(e) > 0$ and $B'(e) > 0$ for all e) and (4) convex and strictly concave respectively ($C''(e) \geq 0$ and $B''(e) < 0$ for all e). Then e^* is a continuous and increasing function with $e^*(0) = 0$. Thus the unique monotone equilibrium has the regime change threshold given by equation (3) of Lemma 3.*

Proof. By standard arguments, e^* is continuous, weakly increasing, and single-valued. Minimum effort $e_{min} = 0$ since 0 is the unique maximizer of $-C(e)$. The proposition is then implied by Lemma 3. $\square$

Clearly, one could give weaker conditions for continuity. The endogenous optimal reward scheme identified in the next section will not satisfy the restrictions on the exogenous reward scheme of Proposition 3. Nonetheless, we will show that the equilibrium effort function is continuous and we will appeal to the characterization of the equilibrium regime threshold of equation (3).

This section provides an analysis of equilibrium with exogenous rewards and punishments. This characterization could be used to address many positive and policy questions. This paper will focus on one: optimal reward schemes.

We conclude this section by reporting a class of examples. Suppose that the exogenous reward scheme is $B(e) = \sqrt{e}$ and the punishment scheme is $C(e) = e^n$, $n \geq 1$. In this case, we have from Proposition 3 that $e^*(p) = (\frac{p}{2n})^{\frac{2}{2n-1}}$ and

$$\theta_{e^*} = \frac{2n-1}{2n+1} (2n)^{\frac{2}{1-2n}}.$$

Figure 1 illustrates.

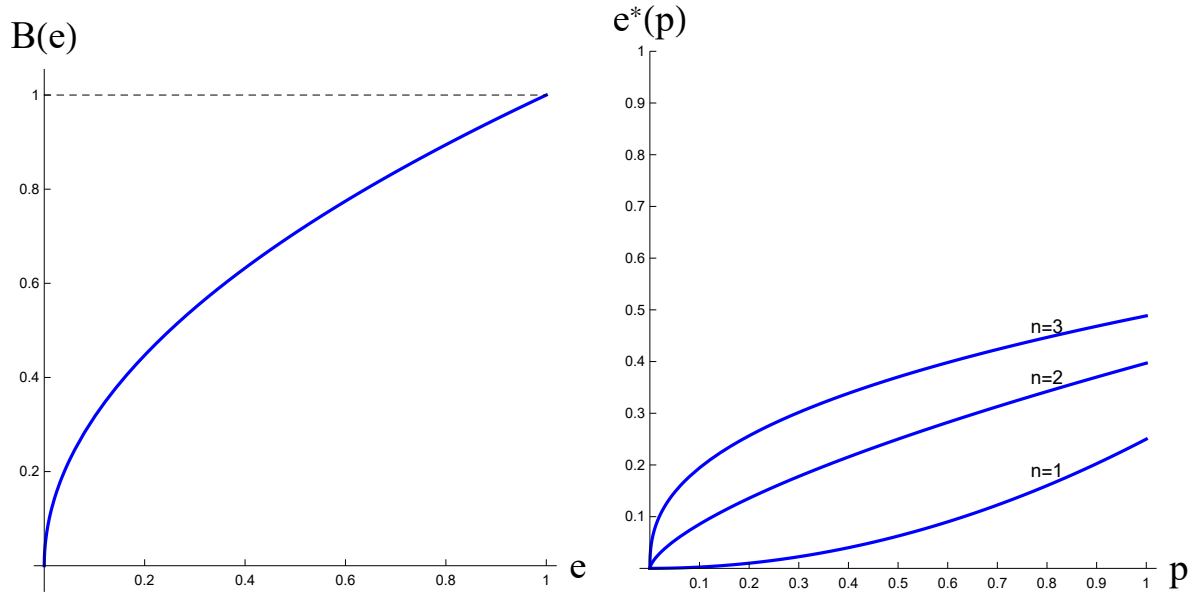


Figure 1: Optimal Effort for Exogenous Rewards. The exogenous reward scheme is $B(e) = \sqrt{e}$ and the exogenous punishment scheme is $C(e) = e^n$, $n \geq 1$. The right panel depicts $e^*(p)$ for $n = 1, 2, 3$.

3 Optimal Reward Schemes

We have analysed the equilibrium effort for fixed benefit and cost functions and identified the implied regime change threshold. We now fix an exogenous cost function C , maintaining the assumptions that C is twice continuously differentiable, strictly increasing and convex, with $C(0) = 0$. These are the properties of the cost function assumed in Proposition 3. We investigate the optimal design of the benefit function $B(e)$ by a revolutionary leader who aims to maximize the likelihood of regime change. Charismatic revolutionary leaders can inspire citizen participation by assigning psychological rewards to different levels of anti-regime activities. Through speeches, writings, and meetings, they activate, create, and manipulate intrinsic motivations by creating identities and innovative framing of events, and by “identification, idealization, and elevation of one or more values presumed basic to prospective constituents” (Snow et al. 1986, p. 469). Unlike material benefits, these psychological rewards are non-rival. For example, in a religious context, if a fighter who is tortured in the righteous struggle to bring down an unjust state or a wicked ruler is to receive one marble castle in heaven, God can build as many castles as there are fighters who deserve them. Still, even charismatic leaders can incite only so much intrinsic motivation in their potential followers. Thus, we require that $B(e) \leq M$ for some exogenous $M > 0$. The upper bound on rewards reflects limitations on the leader’s skills and charisma as well as other exogenous aspects of the environment such as “cultural idioms” (Skocpol 1997) or “repertoire of common symbolics” (Dabashi 1993) that facilitate the leader’s task of creating, activating, and manipulating intrinsic motivations for revolutionary activities.

Proposition 2 implies that $e^*(p)$ is weakly increasing (recall that this proposition did not impose any assumptions on the benefit function). By Proposition 2, we then know that—for any choice of $B(\cdot)$ —there will be a unique equilibrium in which regime change occurs when $\theta \leq \theta_{e^*} = \int_{p=0}^1 e^*(p) dp$. Thus the leader's optimization problem reduces to choosing reward scheme $B(e)$ to maximize the regime change threshold

$$\int_{p=0}^1 e^*(p) dp \quad (8)$$

subject to the optimality of the effort function

$$e^*(p) \in \arg \max_{e \geq 0} p B(e) - C(e),$$

and the feasibility of the reward scheme

$$B(e) \in [0, M] \text{ for all } e. \quad (9)$$

A reward scheme B^* that maximizes this problem is an *optimal reward scheme* (there will be some indeterminacy in optimal reward schemes). An effort function e^* arising in this problem is an *optimal effort function* (there will be an essentially unique optimal effort function). The regime change threshold θ^* that results from the optimal effort function is the (essentially unique) *optimal regime change threshold*. To provide intuition and as a benchmark, we first present the analysis with linear costs, and then present our main result.

3.1 Linear Costs

Proposition 4. *Suppose costs are linear, with constant marginal cost c . An optimal reward scheme B^* is the step function*

$$B^*(e) = \begin{cases} 0 & ; e < M/2c \\ M & ; e \geq M/2c. \end{cases} \quad (10)$$

The optimal effort function e^ is the step function*

$$e^*(p) = \begin{cases} 0 & ; p < 1/2 \\ M/2c & ; p \geq 1/2. \end{cases} \quad (11)$$

And the optimal regime change threshold is

$$\theta_{e^*} = M/4c.$$

Any reward scheme giving rise to the optimal effort function is optimal. Thus, it is enough, for example, that $B^*(0) = 0$, $B^*(M/2c) = M$ and $B^*(e) \leq 2ce$ for $e \leq M/2c$.

To prove the proposition, first assume the optimal reward scheme is a step function, where citizens get the maximum benefit M if they exert at least effort $\hat{e}$; otherwise they get no benefit. In this case, citizens either do not participate (i.e., choose effort 0) or participate with effort $\hat{e}$. A citizen with optimism (probability of regime change) p participates whenever $pM \geq c\hat{e}$, i.e., $p \geq c\hat{e}/M$. Thus, the total effort is $(1 - c\hat{e}/M)\hat{e}$. This is maximized by setting $\hat{e} = M/2c$. Thus we have a proof of the proposition if we could establish that the optimal reward schemes would be a step function.

To complete the proof, we show we can restrict attention to step functions. To do this, we show that this problem reduces to the monopoly pricing problem. Consider a monopolist selling a single unit to buyers whose valuations are uniformly distributed on the interval $[0, M/c]$. The monopolist could sell using a posted price mechanism. At a posted price of $\hat{e}$, buyers with valuations above $\hat{e}$ would buy and pay $\hat{e}$, and buyers with valuations below $\hat{e}$ would not buy. However, the monopolist could also sell using a more complicated mechanism, offering a price schedule for probabilities of being allocated the object. By the revelation principle, we can restrict attention to direct mechanisms. If a type- p buyer has valuation pM/c , a direct mechanism is described by a payment $e^*(p)$ that buyer p will make to the seller and a probability $B(e^*(p))/M$ of receiving the object. The seller's revenue will now be $\int_{p=0}^1 e^*(p)dp$ and incentive compatibility will require that $e^*(p) \in \arg \max_{e \geq 0} p B(e)/c - e$. The revenue is the maximand of our problem and the incentive compatibility condition is the incentive compatibility condition of our problem. But Riley and Zackhauser (1983) established that the optimal mechanism in this problem is a posted price mechanism (see B6rgers (2015, Ch. 2) for a modern textbook treatment), and simple calculations show that the optimal price is $M/2c$. This proves Proposition 4.

Our main result establishes that if the cost function is strictly convex, this step function is “smoothed”: it remains optimal for a mass of the most optimistic citizens to choose a maximum effort level and receive the maximum reward. And it remains optimal (under some conditions) for a mass of the least optimistic citizens to not participate, i.e., choose 0 and receive no benefit. However, in between there are strictly increasing benefits and effort.

We conclude our analysis of the linear case by discussing why bunching at the top starts exactly when $p = 1/2$. Recall that our strategic problem delivers a particular distribution over levels of optimism p : the uniform distribution. In the monopoly problem, this is equivalent to considering a linear demand curve. It is useful to consider what would have happened in the screening problem if p was not uniformly distributed, but distributed according to density f and corresponding cdf F . Under standard assumptions,³ the problem reduces to solving for the critical effort level $\hat{e}$ to receive the (maximum) benefit. As before, the citizen will participate at this critical effort level if $p \geq c\hat{e}/M$. But now total effort would be $(1 - F(c\hat{e}/M))\hat{e}$. Letting e^* be the optimal $\hat{e}$, the first order condition implies

$$e^* = \frac{1 - F(c e^*/M)}{f(c e^*/M)} \frac{M}{c},$$

³In particular, f should satisfy the standard regularity condition that $\frac{1-F(\cdot)}{f(\cdot)}$ is decreasing (decreasing marginal revenue in the monopoly case).

and the corresponding critical optimism level is

$$p^* = ce^*/M = \frac{1 - F(ce^*/M)}{f(ce^*/M)}.$$

When p is distributed uniformly, the above equation implies that $ce^*/M = 1/2$. Thus, $p^* = 1/2$ arises as critical optimism because p is uniformly distributed.

3.2 Strictly Convex Costs and Main Result

Proposition 5. *Suppose costs are strictly convex.*

- For some “maximum effort” $\bar{e} > 0$, the optimal reward scheme is continuous, strictly increasing and strictly convex on the interval $[0, \bar{e}]$ with $B(0) = 0$ and $B(\bar{e}) = M$. Moreover, the marginal benefit of effort $B'(e)$ is strictly greater than the marginal cost of effort $C'(e)$ whenever the marginal benefit is non-zero
- The optimal effort function is continuous and weakly increasing; it is strictly increasing on an interval $[\underline{p}, 1/2]$, equal to 0 when $p \leq \underline{p}$, and equal to $\bar{e}$ when $p \geq 1/2$.
- The critical level of optimism $\underline{p}$ is strictly greater than 0 if and only if the marginal cost is strictly positive at zero effort ($C'(0) > 0$).

To illustrate the proposition, we graph the optimal reward scheme and optimal effort function for some examples. First, suppose $C(e) = e^n$, $n > 1$, and $M = 1$. Figure 2 illustrates the optimal reward and optimal effort schemes in this case for different values of n . The derivation is reported in the appendix. Note that as n approaches 1, we approach the linear case described above (as noted earlier, there is an indeterminacy in the optimal reward scheme, and this limit is piece-wise linear, rather than the step function reported in Proposition 4). This class of examples has $C'(0) = 0$. Figure 3 illustrates an example where $C(e) = e^2 + 0.1e$, so that the marginal cost at 0 is strictly positive, $C'(0) = 0.1$.

Further details about the optimal reward scheme and the optimal effort function, as well as an algorithm for calculating them, are reported below in the proof of the proposition.

The basic features of the optimal reward scheme and optimal effort function are intuitive: both are weakly increasing and continuous, with optimal benefit strictly increasing and ranging from 0 to the maximum. Before presenting a proof, we provide an intuition for further properties. First, consider the shape of the optimal reward scheme when it is strictly increasing. In this case, marginal benefit exceeds the marginal cost. If the marginal cost did exceed the marginal benefit at some effort level, then any effort in the neighborhood of that level could not arise in equilibrium. But then it would be possible to replace the reward scheme with one that was constant in that neighborhood and increasing faster elsewhere, in a way that would increase the overall effort. More mechanically, the first order condition that $pB'(e^*(p)) = C'(e^*(p))$ implies that $B'(e^*(p)) > C'(e^*(p))$ when $p < 1$. To show strict convexity, observe that effort depends on the ratio of marginal costs to marginal

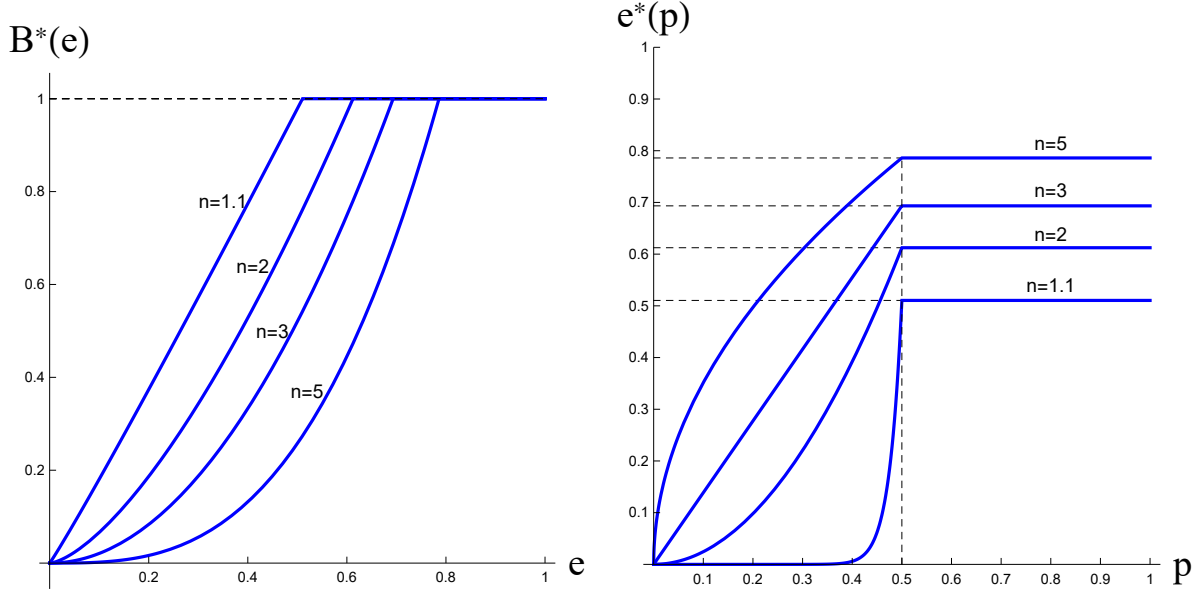


Figure 2: Optimal reward schedule, $B^*(e)$, and its induced effort schedule, $e^*(p)$, for $M = 1$ and cost function $C(e) = e^n$, $n > 1$. As the right panel illustrates, when the cost function approaches the linear $C(e) = e$, effort schedule approaches a step function.

benefits. Increasing marginal benefits will thus have more impact on effort when marginal cost is higher. But marginal cost is itself increasing in effort (by assumption) so marginal benefit will also be increasing. More formally, once we show that $e^*(p)$ is continuous and weakly increasing (and single-valued with $e^*(0) = 0$), then Lemma 3 characterizes θ_{e^*} that the leader seeks to maximize subject to the constraint that $B(e) \leq M$. Differentiating the integrand in θ_{e^*} with respect to $B'(e)$ yields $C'(e)/(B'(e))^2$. Because $C'(e)$ is increasing by assumption, the marginal gains of raising $B'(e)$ are higher for higher effort levels. In contrast, the marginal costs of raising $B'(e)$ to the leader are constant, because the leader has a fixed “budget of slopes” ($\int_{e=0}^{\bar{e}} B'(e)de \leq M$ for some $\bar{e}$). This implies that an optimal $B'(e)$ is increasing, and hence an optimal $B(e)$ is convex.

Now consider the “exclusion” regions where effort is constant. First, consider exclusion at the bottom. Recall that the optimal effort $e^*(p)$ is continuous and increasing, and consider the smallest level of optimism $\underline{p}$ after which $e^*(p)$ becomes strictly positive. Then, $\underline{p} B'(0) = C'(0)$, and hence $C'(0) > 0$ implies $\underline{p} > 0$. In fact, because the marginal costs of raising $B'(e)$ are constant, while its marginal gains are $C'(e)/(B'(e))^2$ as we discussed above, $C'(e)/B'(e)$ is proportional to $\sqrt{C'(e)}$,⁴ and hence $\underline{p} > 0$ exactly when $C'(0) > 0$. Now consider high optimism, and suppose that optimal effort was strictly increasing in a neighborhood around 1. The optimal reward scheme would then have to be strictly increasing in the neighborhood of $\bar{e}$. Now suppose that we considered benefit functions that were equal to M in an open neighborhood of $p = 1$. This would decrease effort on

⁴Formally, $C'(e)/(B'(e))^2 = \lambda \Rightarrow B'(e) = \sqrt{C'(e)/\lambda} \Rightarrow C'(e)/B'(e) = \sqrt{\lambda} \sqrt{C'(e)}$.

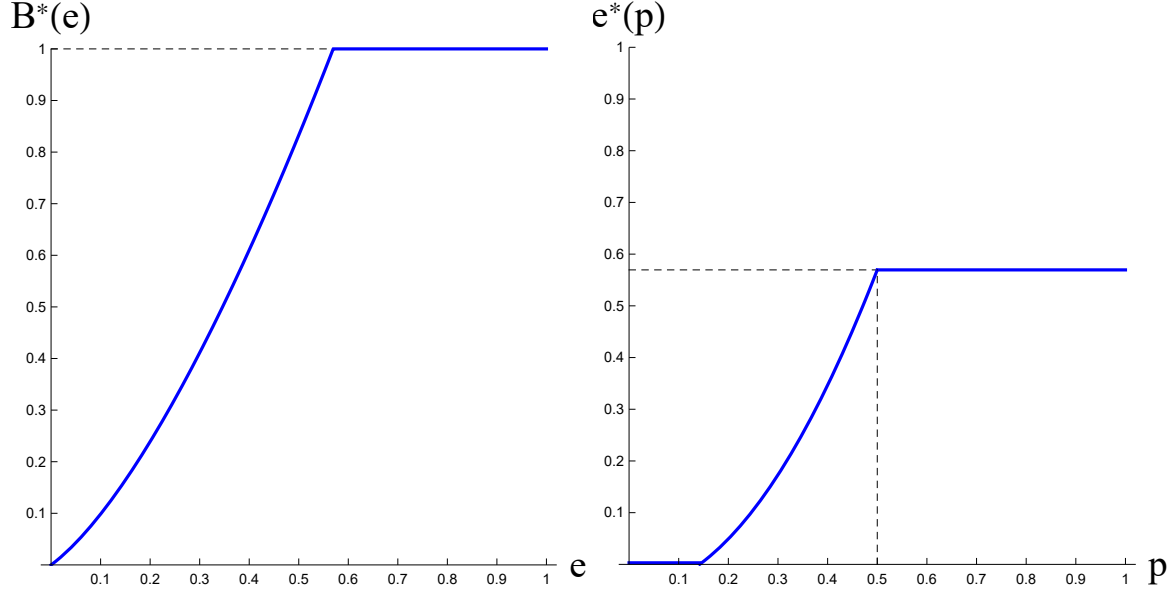


Figure 3: Optimal reward schedule, $B^*(e)$, and its induced effort schedule, $e^*(p)$, for $M = 1$ and cost function $C(e) = e^2 + 0.1e$. As the right panel illustrates, when $C'(0) > 0$, citizens with (endogenously) pessimistic beliefs about the likelihood of regime change do not exert any effort ($\underline{p} > 0$).

the open neighborhood, but would allow optimal effort to be increased at other levels of p . For very small neighborhoods, the latter effect would be of higher order than the former. More formally, the convexity of an optimal $B(e)$ (and that $B(e) \geq C(e)$ in the relevant range) implies that an optimal $B(e)$ hits its constraint M before $C(e)$ does. In turn, this implies that an optimal effort scheme is constant for high values of p .

A last question is why effort reaches its maximum exactly at $p = 1/2$. We already gave an explanation of where $p = 1/2$ comes from in the linear case: it follows from the (endogenous) property that p is uniformly distributed on the interval $[0, 1]$. If p was drawn from a distribution $f(p)$, the corresponding threshold (in the linear cost case) would be when $\frac{1-F(p)}{f(p)} = p$, which corresponds to the point where marginal revenue is zero in the monopoly problem. The same trade-off between intensive and extensive margins arises in the non-linear case. One can show that with a general distribution f of p and under some regularity conditions, the screening problem of Proposition 5 yields the critical p^* that solves

$$\frac{1 - F(p^*)}{f(p^*)} = -\frac{f(p^*)/p^*}{(f(p^*)/p^*)'}.$$

In the uniform case, this expression gives $1 - p^* = -\frac{1/p^*}{-1/(p^*)^2} = p^*$, and hence $p^* = 1/2$.

3.3 Proof

We now present the proof of Proposition 5, formalizing the earlier intuitions. We defer more technical steps to the appendix, where we formalize the leader's problem as a mechanism design problem, and prove that an optimal reward function $e^*(p)$ is continuous and increasing with $e^*(0) = 0$. Then, using Lemma 3, we show that the leader's problem is to choose $(\bar{e}, B')$ to maximize

$$\int_0^{\bar{e}} \left(1 - \frac{C'(e)}{B'(e)}\right) de$$

subject to the marginal benefit constraint that⁵

$$B(\bar{e}) = \int_{e=0}^{\bar{e}} B'(e) de = M.$$

This shows that one can think of the leader's problem as him deciding a maximum level of effort $\bar{e} \geq 0$ to induce, and then deciding how to allocate a fixed supply of marginal benefit to different effort levels $B' : [0, \bar{e}] \rightarrow \mathbb{R}$.

Thus, the leader's problem becomes a constrained point-wise optimization. The Lagrangian is

$$\begin{aligned} L &= \int_0^{\bar{e}} \left(1 - \frac{C'(e)}{B'(e)}\right) de + \lambda \left(M - \int_{e=0}^{\bar{e}} B'(e) de\right) \\ &= \int_0^{\bar{e}} \left(1 - \frac{C'(e)}{B'(e)} - \lambda B'(e)\right) de + \lambda M. \end{aligned}$$

Optimal $B'(e)$ simplifies to a point-wise maximization

$$\frac{\partial L}{\partial B'(e)} \left(1 - \frac{C'(e)}{B'(e)} - \lambda B'(e)\right) = \frac{C'(e)}{[B'(e)]^2} - \lambda = 0. \quad (12)$$

This shows that at the optimum, the marginal gain of raising $B'(e)$, i.e., $C'(e)/[B'(e)]^2$, equals its marginal cost λ . Thus,

$$B'(x) = \frac{1}{\sqrt{\lambda}} \sqrt{C'(x)}, \quad (13)$$

and

$$B(e) = \int_0^e B'(x) dx = \frac{1}{\sqrt{\lambda}} \int_0^e \sqrt{C'(x)} dx. \quad (14)$$

⁵There are two other constraints: $\bar{e} \geq 0$ and $B'(e) \geq C'(e)$. However, as it will be clear from the solution, these two constraints are automatically satisfied.

Combining this with the constraint yields

$$B(\bar{e}) = \frac{1}{\sqrt{\lambda}} \int_0^{\bar{e}} \sqrt{C'(x)} dx = M. \quad (15)$$

Moreover, optimal $\bar{e}$ also satisfies the first order condition

$$\frac{\partial L}{\partial \bar{e}} = 1 - \frac{C'(\bar{e})}{B'(\bar{e})} - \lambda B'(\bar{e}) = 0 \Rightarrow 1 - \frac{C'(\bar{e})}{B'(\bar{e})} = \lambda B'(\bar{e}). \quad (16)$$

Substituting for λ from (12) into (16) yields

$$1 - \frac{C'(\bar{e})}{B'(\bar{e})} = \frac{C'(\bar{e})}{B'(\bar{e})} \Rightarrow \frac{C'(\bar{e})}{B'(\bar{e})} = \frac{1}{2}. \quad (17)$$

So if $\bar{e}$ and λ are the solution to equations (15) and (17), we have solved for B .

From equation (13), we have $\frac{C'(e)}{B'(e)} = \sqrt{\lambda C'(e)} > 0$. This has two implications: (1) because $C(e)$ is strictly convex, $B(e)$ is strictly convex for $0 < e < \bar{e}$; and $\frac{C'(e)}{B'(e)}$ is strictly increasing. Thus, $\frac{C'(\bar{e})}{B'(\bar{e})} = \frac{1}{2}$ implies that $B'(e) \geq C'(e)$. (2) $e^*(p)$ is strictly increasing between $\underline{p} = \sqrt{\lambda} \sqrt{C'(0)}$ and $p = 1/2$; it is equal to 0 for $p \leq \underline{p}$ and equal to $\bar{e}$ for $p \geq 1/2$.

4 Discussion: Inspiring Revolutionary Activities

4.1 Pleasure in Agency

What motivates citizens to participate in revolutions, civil wars, or other anti-regime collective actions, where individual effects are minimal, risks are high, and the fruits of success are public? This question is the essence of Tullock's (1971) "Paradox of Revolution," which specialized Olson's (1965) *Logic of Collective Action* to revolution settings. One possible answer is that participants somehow receive selective material rewards. However, as Blattman and Miguel (2010) argue in their review of the civil war literature, even in such risky environments, "non-material incentives are thought to be common [even] within armed groups. Several studies argue that a leader's charisma, group ideology, or a citizen's satisfaction in pursuit of justice (or vengeance) can also help solve the problem of collective action in rebellion" (p. 15). Similarly, in the broader context of protests and revolutions, researchers have persistently pointed to the critical role of non-material incentives among the participants. For example, based partly on his field work in Eastern Europe, Petersen (2001) argues that non-material status rewards and the feelings of self-respect were essential components of the insurgents' motives; and Pearlman (2016) highlights the role of "moral identity" (e.g., self-respect and "joy of agency") in Syrian protests.

What exactly is the nature of these non-material incentives? In this paper, we focus on the notion of "pleasure in agency." Wood (2003) develops the notion of "pleasure in agency" to capture individuals' motives for participating in contentious collective actions.

Pleasure in agency is “the positive effect associated with self-determination, autonomy, self-esteem, efficacy, and pride that come from the successful assertion of intention” (p. 235). It is “a frequency-based motivation: it depends on the likelihood of success, which in turn increases with the number participating (Schelling 1978; Hardin 1982). Yet the pleasure in agency is undiminished by the fact that one’s own contribution to the likelihood of victory is vanishingly small” (p. 235-6). For example, in her interviews with insurgents during and after the El Salvadoran civil war, she found that insurgents “repeatedly asserted their pride in their wartime activities and consistently claimed *authorship* of the changes that they identified as their work” (p. 231). “They took pride, indeed pleasure, in the successful assertion of their interests and identity...motivated in part by the value they put on being part of the making of history” (p. 18-9).

4.2 The Role of Leaders

Although many scholars agree that “an account of participation in insurgency requires a consideration of the moral and emotional dimensions of participation” (Wood 2003, p. 225), they typically consider these intrinsic motivations as the natural and automatic result of past experiences or the interactions with the regime in periods of contention. We argue that leaders can maximize the likelihood of regime change by strategically creating, activating, or manipulating such feelings and emotions to design an optimal scheme of “pleasure in agency” as the function of citizens’ level of anti-regime activities.

In his review of the literature on revolutions, Goldstone (2001) identifies two distinct types of revolutionary leaders: people-oriented and task-oriented. “People-oriented leaders are those who inspire people, give them a sense of identity and power, and provide a vision of a new and just order” (p. 157). These are leaders who can create, amplify, or transform people’s feelings and identities, e.g., by innovative framing of events and experiences and their personality traits such as their charisma (Snow et al. 1986). In the terminology of the management science organizational behavior, they are “transformational leaders” who have the ability to create inspirational motivations through a variety of psychological mechanisms (Burns 1978, 2003; Bass 1985). In contrast to people-oriented leaders, in Goldstone (2001), “task-oriented leaders are those who can plot a strategy suitable to resources and circumstances” (p. 157), including the repressive capacity of the state and available revolutionary skills, to effectively transform anti-regime sentiments into concrete actions. Combining these separate branches of the literature, we show how a “people-oriented” leader can optimally create and manipulate “pleasure in agency” feelings among potential revolutionaries to maximize the likelihood of regime change. That even the design of optimal “pleasure in agency schemes” requires “a strategy suitable to resources and circumstances” merges the seemingly separate categories of people-oriented and task-oriented leadership.

4.3 Levels of Participation

Charles Tilly defines *Contentious Performances* (Tilly 2008) to be political, collective claim-making actions that fall outside routine social interactions. Such contentious actions take

various forms. Examples include machine-breaking, petitioning, arson (e.g., common in the Swing Rebellion), cattle maiming (e.g., in Tithe War in Ireland), sit-ins in foreign embassies (e.g., in the 1905-07 Iranian Constitutional Revolution), noise-making through cacerolazo (banging on pots and pans, e.g., in 1989 Caracazo in Venezuela), shouting slogans from the roof tops at night (e.g., in the 1979 Iranian Revolution), occupation of a public space or a government building, writing open letters, wearing a particular color (e.g., green wristbands in the 2009 Iranian green movement, orange ribbons in the Ukrainian Orange Revolution, or yellow ribbons in the Yellow Revolution in the Philippines), public meetings, boycotts (e.g., in the American Revolution), parades (e.g., women suffrage movement in the early 20th century), strikes, marches, demonstrations, freedom rides, street blockades, self-immolation, suicide bombing, assassination, hijacking, and guerrilla war.

Some contentious actions require more effort and taking more risks than others. For example, taking up arms against a regime takes far more effort and is far more dangerous than participating in a demonstration, which in turn is typically more costly than participating in a boycott campaign (Shadmehr 2015). That is, different levels of anti-regime activities correspond to different contentious performances. Our analysis shows how a “people-oriented” revolutionary leader, for a given repression structure, $C(e)$, should strategically create and manipulate pleasure in agency, $B(e)$, that citizens obtain from different contentious performances, e , to maximize the likelihood of regime change.

4.4 Professional Revolutionaries

Our main result (Proposition 5) implies that, no matter what the leader’s ability M , the leader’s optimal design of (pleasure in agency) rewards leads to the emergence of a group of citizens who engage in maximum level of anti-regime activities. The creation of this “cadre” is consistent with the notion of professional revolutionaries in Lenin’s treatise, *What Is to be Done?*⁶ Although this group is distinct from other citizens contributing to the cause, there is little difference among the magnitude of duties and the levels of revolutionary activities among the members of the group. Contrasting the organization of workers, who engage in various degrees of contentious activities, with the organization of “revolutionary social-democratic party,” Lenin insisted that “the organization of the revolutionaries must consist first and foremost of people who make revolutionary activity their profession.... In view of this common characteristic of the members of such an organization, all distinctions as between workers and intellectuals, not to speak of distinctions of trade and profession, in both categories, must be effaced” (p. 71). Of course, the members of this group are not identical; some have stronger beliefs in the likelihood of regime change p than others, for example, as a result of their different interpretations of Marxist theory, or differences in mapping from the real world into the theory. However, despite these differences in their beliefs, they constitute a distinct group that exert maximum effort into the revolution—in this sense, they are professional revolutionaries.

⁶It also resonates with the notion of “combat groups” in Mao’s *On Guerrilla Warfare*.

Appendix: Examples and Proofs

Example

Suppose $C(e) = e^n$, $n > 1$, and $M = 1$. What is the optimal $B(e)$? From (14),

$$B(e) = \frac{1}{\sqrt{\lambda}} \int_0^e \sqrt{C'(x)} dx = \frac{1}{\sqrt{\lambda}} \int_0^e \sqrt{n} x^{\frac{n-1}{2}} dx = \frac{1}{\sqrt{\lambda}} \frac{2\sqrt{n}}{n+1} e^{\frac{n+1}{2}}. \quad (18)$$

From (15), $B(\bar{e}) = 1$, which implies

$$\frac{1}{\sqrt{\lambda}} \frac{2\sqrt{n}}{n+1} \bar{e}^{\frac{n+1}{2}} = 1. \quad (19)$$

From (17), $\frac{C'(\bar{e})}{B'(\bar{e})} = \frac{1}{2}$, which implies

$$\frac{n\bar{e}^{n-1}}{\sqrt{\frac{n}{\lambda}} \bar{e}^{\frac{n-1}{2}}} = \sqrt{n\lambda} \bar{e}^{\frac{n-1}{2}} = \frac{1}{2}. \quad (20)$$

From (19), $\sqrt{\lambda} = \frac{2\sqrt{n}}{n+1} \bar{e}^{\frac{n+1}{2}}$. Substituting this $\sqrt{\lambda}$ into equation (20) yields $\frac{2n}{n+1} \bar{e}^n = \frac{1}{2}$, and hence

$$\bar{e} = \left(\frac{n+1}{4n} \right)^{\frac{1}{n}} \quad \text{and} \quad \sqrt{\lambda} = \frac{2\sqrt{n}}{n+1} \left(\frac{n+1}{4n} \right)^{\frac{n+1}{2n}}.$$

Substituting these into equation (18) yields the optimal $B(e) = \left(\frac{4n}{n+1} \right)^{\frac{n+1}{2n}} e^{\frac{n+1}{2}}$ for $e \leq \bar{e}$. Moreover, $B(e) \leq M = 1$ for $e > \bar{e}$. For the purposes of this example, we choose $B(e) = M = 1$ for $e > \bar{e}$. Thus,

$$B^*(e) = \begin{cases} \left(\frac{4n}{n+1} \right)^{\frac{n+1}{2n}} e^{\frac{n+1}{2}} & ; e \leq \bar{e} = \left(\frac{n+1}{4n} \right)^{\frac{1}{n}} \\ M = 1 & ; e > \bar{e}. \end{cases}$$

As expected from our earlier discussion, $B^*(e)$ is *convex* for $e \in [0, \bar{e}]$. Moreover,

$$B^{*'}(e) = \begin{cases} \frac{n+1}{2} \left(\frac{4n}{n+1} \right)^{\frac{n+1}{2n}} e^{\frac{n-1}{2}} & ; e < \bar{e} = \left(\frac{n+1}{4n} \right)^{\frac{1}{n}} \\ 0 & ; e > \bar{e}. \end{cases}$$

How does $e^*(p)$ look like? Recall that $e^*(p) = \arg \max_{e \geq 0} p B^*(e) - e^n$, and hence

$$e^*(p) = \begin{cases} 4^{\frac{1}{n(n-1)}} \left(\frac{n+1}{n} \right)^{\frac{1}{n}} p^{\frac{2}{n-1}} & ; p \in [0, \frac{1}{2}) \\ \bar{e} & ; p \in [\frac{1}{2}, 1]. \end{cases}$$

Figure 2 illustrated this solution.

Proofs

Proof of Proposition 5: As a first step in analyzing this problem, we apply the revelation principle to transform the problem into one where the leader chooses effort levels and benefit levels $\{(e(p), B(p))\}$ depending on the threshold probability p that a citizen assigns to regime change. This mechanism design approach amounts to treating the citizens' endogenous beliefs *as if* they were exogenous and distributed uniformly on $[0, 1]$. However, determining $\{(e(p), B(p))\}$ generates a set of “recommended” revolutionary efforts $E \equiv \{e(p) \text{ s.t. } p \in [0, 1]\}$, and the corresponding rewards $B(e)$ for those $e \in E$. It remains to characterize $B(e)$ for $e \notin E$. The only requirement for such $B(e)$ is that players do not choose it. For example, one could set $B(e) = 0$ for all $e \notin E$. Of course, this choice is not unique. We choose $B(e)$ for $e \notin E$ such that $B(e)$ is constant for $e \notin E$, and $B(e)$ is continuous for all $e \geq 0$. Therefore, we can write the leader's problem as:

$$\begin{aligned} & \max_{\{(e(p), B(p))\}} \int_{p=0}^1 e(p) dp \\ \text{s.t. } & pB(p) - C(e(p)) \geq 0, & \forall p \in [0, 1] \\ & pB(p) - C(e(p)) \geq p B(p') - C(e(p')), & \forall p, p' \in [0, 1] \\ & B(p) \in [0, M], & \forall p \in [0, 1]. \end{aligned}$$

Observe that under this program, the designer can assign any cost to any citizen through their choice of e . To simplify notation, write $h(p) = C(e(p))$.

We can use this new representation of the problem to use standard arguments from screening models. We first establish that $B(p)$ is weakly increasing, and hence $B(p)$ is piecewise continuously differentiable. The incentive compatibility constraints imply $pB(p) - h(p) \geq p B(p') - h(p')$ and $p'B(p') - h(p') \geq p' B(p) - h(p)$. Adding these inequalities implies: $(p - p')[B(p) - B(p')] \geq 0$. Hence, $B(p)$ is weakly increasing, and hence $B(p)$ is piecewise continuously differentiable. Thus, a necessary first order condition is $pB'(p) - h'(p) = 0$ almost everywhere, with the corresponding second order condition $pB''(p) - h''(p) \leq 0$. Differentiating the FOC w.r.t. p yields $B'(p) + pB''(p) - h''(p) = 0$. Thus, the SOC simplifies to $B'(p) \geq 0$. Moreover, because $p B(p) - h(p)$ is increasing in p , the condition $pB(p) - C(e(p)) \geq 0, \forall p \in [0, 1]$, simplifies to $C(e(0)) = h(0) = 0$. Further, because $B'(p) \geq 0$, the constraint $B(p) \in [0, M], \forall p \in [0, 1]$, can be replaced by $B(0) \geq 0$ and $B(1) \leq M$. Combining these results, the leader's problem becomes:

$$\max_{\{(e(p), B(p))\}} \int_{p=0}^1 e(p) dp \tag{21}$$

$$pB'(p) - h'(p) = 0, \quad h(0) = 0 \tag{22}$$

$$B'(p) \geq 0 \tag{23}$$

$$B(0) \geq 0, \quad B(1) \leq M. \tag{24}$$

We can re-write this problem, letting $\Pi(\cdot) = C^{-1}(\cdot)$, so that $\Pi(h(p)) = e(p)$. Then, the

leader's problem (21)-(24) becomes:

$$\max_{\{(B(p), h(p))\}} \int_{p=0}^1 \Pi(h(p)) dp \quad (25)$$

$$h'(p) = pB'(p), \quad h(0) = 0 \quad (26)$$

$$B'(p) \geq 0, \quad B(0) \geq 0, \quad \text{and} \quad B(1) \leq M. \quad (27)$$

Because incentives are created by the slope of benefits, it is clear that $B(0) = 0$ and $B(1) = M$. To proceed, we use the optimal control techniques by defining two state variables. Taking a similar approach to Kamien and Schwartz (2012, p. 244-6), let (h, B) be the state and B' be the control, so that $h' = pB'$. The Hamiltonian and Lagrangian are:

$$H = \Pi(h) + \lambda_h pB' + \lambda_B B'.$$

$$L = H + \mu B'.$$

Then, by the maximum principle,

$$\frac{\partial L}{\partial B'} = \lambda_B + \lambda_h p + \mu = 0, \quad \mu \geq 0, \quad \mu B' = 0. \quad (28)$$

$$\lambda'_B(p) = -\frac{\partial L}{\partial B} = 0. \quad (29)$$

$$\lambda'_h(p) = -\frac{\partial L}{\partial h} = -\Pi'(h(p)), \quad \lambda_h(1) = 0. \quad (30)$$

$$B' \quad \text{must maximize } H \text{ given optimal } (B, h) \text{ and } (\lambda_B, \lambda_h). \quad (31)$$

Moreover, we recognize that $B(1) = M$ is captured by having a free $\lambda_B(1)$ as opposed to $\lambda_h(1) = 0$ that stems from a free $h(1)$. Now, because $C(e)$ is strictly convex, its inverse $\Pi(\cdot)$ is strictly concave. This allows us to apply Arrow's (1966) approach together with a theorem from Seierstad and Sydsæter (1993) to show that $h(p)$ and $B(p)$ have no jumps.

Lemma 4. *Optimal $h(p)$ and $B(p)$ have no jump.*

Proof of Lemma 4: We show that an optimal $h(p)$ has no jump, which then implies that an optimal $B(p)$ has no jump. First, we prove that there is no interior jumps at any $p \in (0, 1)$. Our proof is based on Arrow (1966, p. 11-3). Suppose $h(p)$ has a jump at $\tau \in (0, 1)$. Let $h(\tau^+) \equiv \lim_{p \rightarrow \tau^+} h(p)$ and $h(\tau^-) \equiv \lim_{p \rightarrow \tau^-} h(p)$. Because $h(p)$ is increasing, $h(\tau^-) < h(\tau^+)$. From (29) and (30), recall that $\lambda'_B(p) + p\lambda'_h(p) = -p\Pi'(h(p))$ in $(\tau - \epsilon, \tau) \cup (\tau, \tau + \epsilon)$ for some $\epsilon > 0$. Because $\Pi(x)$ is strictly concave, $\Pi'(h(\tau^+)) < \Pi'(h(\tau^-))$, and hence:

$$\lambda'_B(\tau^-) + \tau \lambda'_h(\tau^-) < \lambda'_B(\tau^+) + \tau \lambda'_h(\tau^+). \quad (32)$$

Moreover, from (28), $\lambda_B(p) + p\lambda_h(p) \leq 0$ and $\lambda_B(\tau) + \tau\lambda_h(\tau) = 0$.⁷ Further, from condition (74) of Theorem 7 in Seierstad and Sydsæter (1993, p. 197), $\lambda_h(p)$ and $\lambda_B(p)$ are continuous

⁷More formally, $\lambda_B(\tau) + \tau\lambda_h(\tau) = 0$ obtains from conditions (74) and (75) in Theorem 7 of Seierstad and Sydsæter (1993, p. 197).

at τ . Thus, in a right neighborhood of τ , we have $\lambda'_B(p) + p\lambda'_h(p) + \lambda_h(p) \leq 0$, and in a left neighborhood of τ , we have $\lambda'_B(p) + p\lambda'_h(p) + \lambda_h(p) \geq 0$. But this contradicts (32).

Next, we prove that there is no jump at $p = 0$ or $p = 1$. From (29), $\lambda_B(p) = \text{constant} \equiv \bar{\lambda}$ on $p \in (0, 1)$. Moreover, $\lambda_B(p)$ is continuous (Seierstad and Sydsæter 1993, p. 197). Thus, $\lambda_B(p) = \bar{\lambda}$, for $p \in [0, 1]$. From conditions (74) and (75) in Theorem 7 of Seierstad and Sydsæter (1993, p. 197), $\lambda_B(\tau) + \tau\lambda_h(\tau) = 0$. Combining this with $\lambda_B(p) = \bar{\lambda}$ and (30), $\bar{\lambda} + \tau \int_{x=\tau}^1 \Pi'(h(x))dx = 0$. In particular, if $\tau = 0$ or $\tau = 1$, then $\bar{\lambda} = 0$. But from (28), $\bar{\lambda} + \lambda_h p \leq 0$ and $\lambda_h(p) = \int_{x=p}^1 \Pi'(h(x))dx$ is positive for some $p \in (0, 1)$. Thus, $\bar{\lambda} < 0$. A contradiction. $\square$

From (29) and (30),

$$\lambda_B(p) = \text{constant} = \bar{\lambda} \quad \text{and} \quad \lambda_h(p) = - \int_{x=p}^1 \lambda'_h(x)dx = \int_{x=p}^1 \Pi'(h(x))dx. \quad (33)$$

Moreover, from (28), $\bar{\lambda} + p \lambda_h(p) \leq 0$. Because at $p \in (0, 1)$, $p \lambda_h(p) = p \int_{x=p}^1 \Pi'(h(x))dx > 0$, we must have $\bar{\lambda} < 0$.

Lemma 5. *There is no interior bunching. That is, there is no $0 < p_1 < p_2 < 1$ such that $h'(p) = 0$ for $p \in (p_1, p_2)$, with $h'(p) > 0$ in a left neighborhood of p_1 and in a right neighborhood of p_2 .*

Proof of Lemma 5: Suppose not. Let $h(p) = \bar{h}$ for $p \in [p_1, p_2]$. Because $h(p)$ and hence $B(p)$ must be strictly increasing in a left neighborhood of p_1 and in a right neighborhood of p_2 , $\mu = 0$ in those neighborhoods. Moreover, $\lambda_B(p)$ and $\lambda_h(p)$ are continuous. Therefore, (28) implies that $\lambda_B(p_1) + p_1\lambda_h(p_1) = 0 = \lambda_B(p_2) + p_2\lambda_h(p_2)$. Because $\lambda_B(p) = \bar{\lambda}$ from (33), we have,

$$p_1\lambda_h(p_1) = p_2\lambda_h(p_2). \quad (34)$$

Further, from (33),

$$\lambda_h(p_1) = \int_{p_1}^{p_2} \Pi'(h(x))dx + \int_{p_2}^1 \Pi'(h(x))dx = (p_2 - p_1) \Pi'(\bar{h}) + \lambda_h(p_2). \quad (35)$$

Combining equations (34) and (35) yields:

$$p_1 \Pi'(\bar{h}) = \lambda_h(p_2). \quad (36)$$

Next, consider $p_3 \in (p_1, p_2)$. From (28), $\lambda_B(p_3) + p_3\lambda_h(p_3) \leq 0 = \lambda_B(p_2) + p_2\lambda_h(p_2)$. Hence,

$$p_3\lambda_h(p_3) \leq p_2\lambda_h(p_2). \quad (37)$$

Mirroring the calculations of equation (35), we have:

$$\lambda_h(p_3) = (p_2 - p_3) \Pi'(\bar{h}) + \lambda_h(p_2). \quad (38)$$

Combining equations (37) and (38) yields:

$$p_3 \Pi'(\bar{h}) \leq \lambda_h(p_2). \quad (39)$$

From equations (36) and (39), $p_3 \Pi'(\bar{h}) \leq p_1 \Pi'(\bar{h})$. However, because $p_3 > p_1$ and $\Pi'(\bar{h}) > 0$, we must have $p_3 \Pi'(\bar{h}) > p_1 \Pi'(\bar{h})$, a contradiction. $\square$

Next, we solve the problem ignoring the constraint $B'(p) \geq 0$, and subsequently check whether and when this constraint binds. Without $B'(p) \geq 0$, $\mu = 0$, and hence from (28) and (33),

$$p \int_{x=p}^1 \Pi'(h(x)) dx + \bar{\lambda} = 0, \text{ and hence } \frac{d}{dp} \left\{ p \int_{x=p}^1 \Pi'(h(x)) dx + \bar{\lambda} \right\} = 0 \quad (40)$$

Differentiating (40) with respect to p yields $\Pi'(h(p)) = -\frac{\bar{\lambda}}{p^2}$, and hence

$$h(p) = \Pi'^{-1} \left(-\frac{\bar{\lambda}}{p^2} \right). \quad (41)$$

Moreover, Because $\Pi(h)$ is strictly concave, $h(p)$ is strictly increasing, and hence $B(p)$ is strictly increasing as far as the constraints $B(0) = 0$ and $B(1) = M$ are satisfied. Therefore, an optimal $B(p)$ takes the following form:

$$B(p) = \begin{cases} 0 & ; p \in [0, p_1] \\ \text{strictly increasing function} & ; p \in [p_1, p_2] \\ M & ; p \in [p_2, 1], \end{cases}$$

for $0 \leq p_1 < p_2 \leq 1$. This, in turn, implies a similar form for $h(p)$, and hence for optimal effort schedule:

$$e^*(p) = \begin{cases} 0 & ; p \in [0, \underline{p}] \\ \text{strictly increasing and satisfies the first order condition} & ; p \in [\underline{p}, \bar{p}] \\ \bar{e} & ; p \in [\bar{p}, 1], \end{cases} \quad (42)$$

for some $0 \leq \underline{p} < \bar{p} \leq 1$ and $\bar{e} > 0$.

Equation (42) together with Lemma 3 allows us to write the leader's objective function as:

$$\theta^* = \int_{e=0}^{\bar{e}} \left(1 - \frac{C'(e)}{B'(e)} \right) de,$$

where we recognize that $e^*(p)$ satisfying the first order condition implies $B'(e) \geq C'(e)$, and that $\underline{p} = \frac{C'(0)}{B'(0)}$ and $\bar{p} = \frac{C'(\bar{e})}{B'(\bar{e})}$. Thus, we can formulate the leader's problem as:

$$\begin{aligned} \max_{B', \bar{e}} \quad & \int_0^{\bar{e}} \left(1 - \frac{C'(e)}{B'(e)} \right) de \\ \text{s.t.} \quad & B(\bar{e}) = M, \ B(0) = 0, \ B'(e) \geq C'(e), \ \bar{e} \geq 0, \end{aligned} \quad (43)$$

Writing $B(\bar{e})$ as $B(\bar{e}) = \int_{e=0}^{\bar{e}} B'(e)de$, (43) can be written as:

$$\begin{aligned} \max_{B', \bar{e}} \quad & \int_0^{\bar{e}} \left(1 - \frac{C'(e)}{B'(e)}\right) de \\ \text{s.t.} \quad & \int_{e=0}^{\bar{e}} B'(e)de = M, \quad B(0) = 0, \quad B'(e) \geq C'(e), \quad \bar{e} \geq 0. \end{aligned}$$

This is the maximization problem that we analyze in Section 3.3 of the text. The rest of the proof is in the text. $\square$

5 References

- Arrow, Kenneth. 1966. "Optimal Capital Policy with Irreversible Investment." Institute for Mathematical Studies in Social Science, Stanford University.
- Baron, David, and Roger Myerson. 1982. "Regulating a Monopolist with Unknown Costs." *Econometrica* 50 (4): 911-30.
- Bass, Bernard. 1985. *Leadership and Performance Beyond Expectations*. NY: Free Press.
- Blattman, Christopher, and Edward Miguel. 2010. "Civil War." *Journal of Economic Literature* 48: 3-57.
- Börgers, Tilman. 2015. *An Introduction to the Theory of Mechanism Design*. New York: Oxford University Press.
- Bueno de Mesquita, Ethan. 2010. "Regime Change and Revolutionary Entrepreneurs." *American Political Science Review* 104 (3): 446-66.
- Burns, James MacGregor. 1978. *Leadership*. New York: Harper & Row Publishers.
- Burns, James MacGregor. 2003. *Transforming Leadership: A New Pursuit of Happiness*. New York: Grove Press.
- Chen, Heng, and Wing Suen. 2016. "Falling Dominoes: A Theory of Rare Events and Crisis Contagion." *American Economic Journal: Microeconomics* 8: 1-29.
- Dabashi, Hamid. 1993. *Theology of Discontent: The Ideological Foundation of the Islamic Revolution in Iran*. New York: New York University Press.
- Edmond, Chris. 2013. "Information Manipulation, Coordination and Regime Change." *Review of Economic Studies* 80: 1422-58.
- Egorov, Georgy, and Konstantin Sonin. 2017. "Elections in Non-Democracies." Mimeo.
- Frankel, David, Ady Pauzner and Stephen Morris. 2003. "Equilibrium Selection in Global Games with Strategic Complementarities," *Journal of Economic Theory* 108: 1-44.
- Goldstone, Jack A. 2001. "Toward a Forth Generation of Revolutionary Theory." *Annual Review of Political Science* 4: 139-87.
- Guesnerie, Roger, and Jean-Jacques Laffont. 1984. "A Complete Solution to a Class of Principal-Agent Problems with an Application in the Control of a Self-Managed Firm." *Journal of Public Economics* 25: 329-69.
- Guimaraes, Bernardo, and Stephen Morris. 2007. "Risk and Wealth in a Model of Self-fulfilling Currency Attack." *Journal of Monetary Economics* 54: 2205-30.

- Hardin, Russell. 1982. *Collective Action*. Baltimore, MD: Johns Hopkins University Press.
- Hartline, Jason D. 2016. *Mechanism Design and Approximation*. Online Manuscript.
- Kamien, Morton, and Nancy Schwartz. 2012. *Dynamic Optimization: The Calculus of Variations and Optimal Control in Economics and Management (2nd Edition)*. Dover.
- Laffont, Jean-Jacques, and Jacques Robert. 1996. "Optimal Auction with Financially Constrained Buyers." *Economics Letters* 52: 181-86.
- Maskin, Eric, and John Riley. 1984. "Monopoly with Incomplete Information." *RAND Journal of Economics* 15 (2): 171-96.
- Mookherjee, Dilip, and I. P. L. Png. 1994. "Marginal Deterrence in Enforcement of Law." *Journal of Political Economy* 102 (5): 1039-66.
- Morris, Stephen, and Hyun Song Shin. 2003. "Global Games: Theory and Application." In *Advances in Economics and Econometrics, Theory and Applications, Eighth World Congress, Volume I*, edited by M. Dewatripont, L. P. Hansen, and S. J. Turnovsky. New York: Cambridge University Press.
- Mussa, Michael, and Sherwin Rosen. 1978. "Monopoly and Product Quality." *Journal of Economic Theory* 18: 301-17.
- Olson, Mancur. 1965. *The Logic of Collective Action: Public Goods and the Theory of Groups*. Cambridge, MA: Harvard University Press.
- Pearlman, Wendy. 2016. "Moral Identity and Protest Cascades in Syria." *British Journal of Political Science*. Forthcoming.
- Petersen, Roger D. 2001. *Resistance and Rebellion: Lessons from Eastern Europe*. New York: Cambridge University Press.
- Quigley, Daniel, and Frederik Toscani. 2016. "The roles of transparency in regime change: Striking when the iron's gone cold." Mimeo.
- Riley, John, and Richard Zackhauser. 1983. "Optimal Selling Strategies: When to Haggle, When to Hold Firm." *Quarterly Journal of Economic* 98 (2): 267-289.
- Schelling, Thomas C. 1978. *Micromotives and Macrobbehavior*. New York: Norton.
- Seierstad, Atle, and Knut Sydsæter. 1993. *Optimal Control Theory with Economic Applications*. Amsterdam, The Netherlands: Elsevier.
- Shadmehr, Mehdi. 2015. "Extremism in Revolutionary Movements." *Games and Economic Behavior* 94: 97-121.
- Shadmehr, Mehdi. 2017. "Capitalists in Revolution." Mimeo.
- Shen, Lin, and Junyuan Zou. 2016. "Intervention with Voluntary Participation in Global Games." Mimeo.
- Skocpol, Theda. 1997. "Cultural Idioms and Political Ideologies in the Revolutionary Reconstruction of State Power: A Rejoinder to Sewell." In *Social Revolutions in the Modern World*. Ed. by Theda Skocpol. New York: Cambridge University Press.
- Snow, David, Burke Rochford, Jr., Steven Worden, and Robert Benford. 1986. "Belief Alignment Processes, Micromobilization, and Movement Participation." *American Sociological Review* 51: 464-81.
- Tilly, Charles. 2008. *Contentious Performances*. New York: Cambridge University Press.
- Tullock, Gordon. 1971. "The Paradox of Revolution." *Public Choice* 11: 89-99.

- Tyson, Scott, and Alastair Smith. 2016. "Two-Sided Coordination: Competition in Collective Action." *Journal of Politics*. *Forthcoming*.
- Wood, Elisabeth J. 2003. *Insurgent Collective Action and Civil War in El Salvador*. New York, NY: Cambridge University Press.