

Online Appendix

LEMMA A. *Increasing repression causes the probability of a reassertion of power to increase in the anocratic equilibrium if, and only if, condition (P) is not satisfied. This probability is always increasing in the despotic equilibrium.* $\square$

Proof of Lemma A The probability of reassertion of power is just the probability of neither citizen being actively opposed, $(1 - \lambda_A)^2$ in the anocratic equilibrium, and $(1 - \lambda_D)^2$ in the despotic equilibrium. Thus, its behavior is the inverse of λ_A and λ_D , respectively. The claim follows immediately from Lemma 5 for the anocratic equilibrium, and (4) for the despotic one. $\blacksquare$

LEMMA B. *If (P) is not satisfied, the probability of a costly civil conflict is decreasing in repression in the anocratic equilibrium. If (P) is satisfied, then it is decreasing if, and only if,*

$$1 + \sqrt{3} \geq \left[3 \left(\frac{\zeta}{\pi} - 1 \right) + \sqrt{3} \right] \bar{w},$$

otherwise it is concave (increasing for low values of k , and then decreasing). In the despotic equilibrium, the probability is always zero. $\square$

Proof of Lemma B. For civil conflict to occur, both dissidents and regime supporters have to be active, for which the probability is $2\lambda_A\varphi_A$, so:

$$\frac{d \text{ Conflict}}{d k} = 2 \left(\varphi_A \cdot \frac{d \lambda_A}{d k} + \lambda_A \cdot \frac{d \varphi_A}{d k} \right) \geq 0.$$

Since $\frac{d \varphi_A}{d k} < 0$ by Lemma 4, if $\frac{d \lambda_A}{d k} \leq 0$, that is, (P) does not hold, then this derivative is negative,

which establishes the first part of the claim. Suppose now that (P) obtains, so $\frac{d\lambda_A}{dk} > 0$. From the proof of Lemma 6, we can rewrite the derivative

$$(\zeta - 4\pi\lambda_A)\varphi_A + \left[4(\lambda_A + \pi\varphi_A) - 3 - \overline{w}\right]\lambda_A \geq 0,$$

which we can simplify to

$$\zeta\varphi_A \geq (3 - 4\lambda_A + \overline{w})\lambda_A.$$

Substituting (6) into (5) and simplifying yields

$$\zeta\varphi_A = 1 - 2k - (3 - 2\lambda_A - \overline{w})\lambda_A,$$

which means that we need to determine

$$1 - 2k - (3 - 2\lambda_A - \overline{w})\lambda_A \geq (3 - 4\lambda_A + \overline{w})\lambda_A,$$

which simplifies to

$$\frac{1 - 2k}{6} \geq (1 - \lambda_A)\lambda_A.$$

Observe now that the left-hand side is decreasing in k while the right-hand side is increasing (because $\lambda_A < 1/2$ means that it is increasing in λ_A , and λ_A is increasing in k by our supposition), we conclude that the sign can change at most once. Moreover, since

$$\lim_{k \rightarrow k^*} \frac{1 - 2k}{6} < \lim_{k \rightarrow k^*} (1 - \lambda_A)\lambda_A = (1 - \lambda_D)\lambda_D \Leftrightarrow 0 < 1 + 2k^*(4 - k^*),$$

it follows that for high enough k , the probability of conflict is decreasing. But this and the fact that the sign can change at most once imply that there are only two possibilities: either this probability is always decreasing or it is increasing for some $k \in (0, \hat{k})$ and decreasing for $k \in (\hat{k}, k^*)$. This probability can be strictly decreasing if, and only if,

$$\lim_{k \rightarrow 0} \frac{1-2k}{6} = \frac{1}{6} \leq \lim_{k \rightarrow 0} (1 - \lambda_A) \lambda_A \Leftrightarrow \lim_{k \rightarrow 0} \lambda_A \geq \frac{1 - \sqrt{1/3}}{2}.$$

Since (6) tells us that

$$\lim_{k \rightarrow 0} \varphi_A = \frac{\overline{w}}{2\pi},$$

we can use (5) to obtain the quadratic in the limit as $k \rightarrow 0$:

$$-2\lambda_A^2 + (3 - \overline{w})\lambda_A - \left(1 - \frac{\zeta\overline{w}}{2\pi}\right) = 0,$$

whose discriminant is

$$(3 - \overline{w})^2 - 8 \left(1 - \frac{\zeta\overline{w}}{2\pi}\right) > 0.$$

Since the larger root exceeds $1/2$, the only admissible solution is

$$\lim_{k \rightarrow 0} \lambda_A = \frac{3 - \overline{w} - \sqrt{(3 - \overline{w})^2 - 8 \left(1 - \frac{\zeta\overline{w}}{2\pi}\right)}}{4}$$

Thus, the necessary and sufficient condition for the probability of conflict to be decreasing is

$$3 - \overline{w} - \sqrt{(3 - \overline{w})^2 - 8 \left(1 - \frac{\zeta\overline{w}}{2\pi}\right)} \geq 2 \left(1 - \sqrt{1/3}\right),$$

which simplifies to the condition stated in the lemma. If this condition is not satisfied, then the probability must be concave. ■

LEMMA C. *Repression causes the probability of a velvet revolution to increase in the anocratic equilibrium and decrease in the despotic equilibrium.* □

Proof of Lemma C. The probability of a velvet revolution (only regime opponents are active with positive probability) in the anocratic equilibrium is $\lambda_A^2 + 2\lambda_A(1 - \lambda_A - \varphi_A) = 2\lambda_A - \lambda_A^2 - 2\lambda_A\varphi_A$, so we need to show that

$$\frac{d \text{VR}}{d k} = 2 \left[(1 - \lambda_A - \varphi_A) \cdot \frac{d \lambda_A}{d k} - \lambda_A \cdot \frac{d \varphi_A}{d k} \right] > 0.$$

Since $\frac{d \varphi_A}{d k} < 0$ (Lemma 4), the inequality obtains whenever $\frac{d \lambda_A}{d k} \geq 0$. We now establish that it also does when $\frac{d \lambda_A}{d k} < 0$. Recall from the proof of Lemma 6 that

$$\frac{d \Omega_A}{d k} = 2 \left[\pi \lambda_A \cdot \frac{d \varphi_A}{d k} - (1 - \lambda_A - \pi \varphi_A) \cdot \frac{d \lambda_A}{d k} \right] < 0.$$

But now we obtain

$$\lambda_A \cdot \frac{d \varphi_A}{d k} < \pi \lambda_A \cdot \frac{d \varphi_A}{d k} < (1 - \lambda_A - \pi \varphi_A) \cdot \frac{d \lambda_A}{d k} < (1 - \lambda_A - \varphi_A) \cdot \frac{d \lambda_A}{d k},$$

where the first inequality follows from $\frac{d \varphi_A}{d k} < 0$, the second from $\frac{d \Omega_A}{d k} < 0$ above, and the third from our supposition that $\frac{d \lambda_A}{d k} < 0$.

In the despotic equilibrium, the probability of a velvet revolution is just $\lambda_D^2 + 2\lambda_D(1 - \lambda_D)$, which

means that

$$\frac{d \text{VR}}{d k} = 2(1 - \lambda_D) \cdot \frac{d \lambda_D}{d k} < 0,$$

where the inequality follows from (4). ■

LEMMA D. *If $\pi > 1/2$ then (P) is monotonic in θ : there exists $\tilde{\theta}$ such that it holds if, and only if, $\theta > \tilde{\theta}$.* □

Proof. Taking the derivative of the left-hand side with respect to θ yields:

$$1 + \frac{4}{\sqrt{1 + 8k^*}} \cdot \frac{d k^*}{d \theta} > 0,$$

where we establish the inequality as follows. Since

$$\frac{d h}{d \theta} = \frac{(1 - \pi)h(\bar{w})}{\sqrt{(3 + \bar{w})^2 - 8}},$$

we obtain:

$$\frac{d k^*}{d \theta} = \bar{w} \cdot \frac{d h}{d \theta} - (1 - \pi)h(\bar{w}) = (1 - \pi)h(\bar{w}) \left[\frac{\bar{w}}{\sqrt{(3 + \bar{w})^2 - 8}} - 1 \right] < 0,$$

where the inequality follows from the fact that $\bar{w} < \sqrt{(3 + \bar{w})^2 - 8}$. We thus need to show that

$$4(1 - \pi)h(\bar{w}) \left[1 - \frac{\bar{w}}{\sqrt{(3 + \bar{w})^2 - 8}} \right] < \sqrt{1 + 8\bar{w}h(\bar{w})}. \quad (16)$$

We first show that the left-hand side is decreasing in $\bar{w}$. We can rewrite it as

$$4(1 - \pi) \left[\frac{h(\bar{w})}{\sqrt{(3 + \bar{w})^2 - 8}} \right] \left[\sqrt{(3 + \bar{w})^2 - 8} - \bar{w} \right],$$

and we note that since $h(\bar{w})$ is decreasing,

$$\frac{dh}{d\bar{w}} = \left(\frac{1}{4} \right) \left[1 - \frac{3 + \bar{w}}{\sqrt{(3 + \bar{w})^2 - 8}} \right] < 0,$$

the first bracketed term is decreasing. It suffices to show that so does the second bracketed term.

Taking the derivative with respect to $\bar{w}$ yields

$$(1 - \pi) \left[1 - \frac{3 + \bar{w}}{\sqrt{(3 + \bar{w})^2 - 8}} \right] = 4(1 - \pi) \cdot \frac{dh}{d\bar{w}} < 0,$$

which holds. Since $\bar{w} h(\bar{w})$ is increasing, it will be sufficient to establish (16) as $\bar{w} \rightarrow 0$. But then

(16) reduces to $2(1 - \pi) < 1$, which holds under the assumption that $\pi > 1/2$. ■

complete-info

Why Not Rely On Retaliatory Repression?

While preventive repression can be effective whenever the ruler can implement it at sufficiently high levels, it is distinctly inimical to the ruler's survival when he cannot. Perhaps he could do better with retaliatory repression? After all, unlike preventive repression, which penalizes any political action irrespective of its content or consequences, retaliatory repression imposes costs only when conflict actually occurs, and then only on the side that happens to lose it.²⁷

We now show that retaliatory repression is less useful as a policy tool for the ruler than preventive repression. We first establish the analogue to Lemma 4: retaliatory repression also deters supporters from taking action.

LEMMA E. *Increasing retaliatory repression makes regime supporters less likely to be active in the anocratic equilibrium.* □

Proof of Lemma E. Consider the anocratic equilibrium. Since both (5) and (6) must hold in equilibrium, we differentiate both their sides with respect to θ :

$$\left(3 - 4\lambda_A - 2\pi\varphi_A\right) \cdot \frac{d\lambda_A}{d\theta} + \pi\varphi_A = -\left(\zeta - 2\pi\lambda_A\right) \cdot \frac{d\varphi_A}{d\theta} \quad (17)$$

$$-\left(\bar{w} - 2\pi\varphi_A\right) \cdot \frac{d\lambda_A}{d\theta} + (1 - \pi)\lambda_A = -2\pi\lambda_A \cdot \frac{d\varphi_A}{d\theta} \quad (18)$$

27. It is important to bear in mind that the model assumes that the opponents cannot credibly commit not to punish the supporters if the ruler is toppled. This assumption is fairly realistic when the new ruler is another authoritarian but it is also not out of the question if the new regime is a (transitional) democracy. In these contexts, making the alternative assumption that only regime opponents suffer from retaliatory repression seems less justified and more demanding.

Since $3 - 4\lambda_A - 2\pi\varphi_A > 0$ and $\zeta - 2\pi\lambda_A > 0$, (17) implies that

$$\frac{d\lambda_A}{d\theta} \geq 0 \Rightarrow \frac{d\varphi_A}{d\theta} < 0,$$

and since $\bar{w} - 2\pi\varphi_A > 0$, (18) implies that

$$\frac{d\lambda_A}{d\theta} \leq 0 \Rightarrow \frac{d\varphi_A}{d\theta} < 0.$$

Since $\frac{d\varphi_A}{d\theta} < 0$ must obtain in every possible case, the claim holds. ■

The intuition is simple: the more repressive the regime is to its opponents when they lose, the more its supporters fear what will happen to them if they lose. Although mediated through the probability of winning, the effect on supporters is analogous to that of preventive repression.

The effect on regime opponents is a bit more complicated because it turns out that λ_A might not be monotonic in θ as it is in k . It is possible for some relatively modest retaliatory repression can cause opponents to be less likely to act in the anocratic equilibrium. However, this deterrent effect is quickly outweighed by the incentive to act provided by regime supporters dropping out at even higher rates. This makes retaliatory repression relative unattractive to the ruler in the anocratic equilibrium except perhaps at very low levels, as the following result shows.

LEMMA F. *Increasing retaliatory repression in the anocratic equilibrium might initially cause regime opponents to be less likely to act, but always makes them more likely to do so once the penalties become sufficiently severe. Nevertheless, the probability that opponents act is always smaller in the anocratic equilibrium than in the despotic one (where it is constant): $\lambda_A < \lambda_D$.* □

Proof of Lemma F. We need to show that λ_A is convex. We can simplify (17) and (18) to obtain:

$$\begin{aligned}\gamma \cdot \frac{d\lambda_A}{d\theta} &= \lambda_A \left[(1-\pi)\zeta - 2\pi((1-\pi)\lambda_A + \pi\varphi_A) \right] \equiv \lambda_A f(\theta) \\ \gamma \cdot \frac{d\varphi_A}{d\theta} &= -(\bar{w} - 2\pi\varphi_A)\pi\varphi_A - (1-\pi)(3 - 4\lambda_A - 2\pi\varphi_A)\lambda_A,\end{aligned}$$

where $\gamma \equiv (\bar{w} - 2\pi\varphi_A)\zeta + 2\pi\lambda_A(3 - 2\bar{w} - 4\lambda_A) > 0$. This tells us that

$$\text{sgn}\left(\frac{d\lambda_A}{d\theta}\right) = \text{sgn}(f(\theta)) \quad \text{and} \quad \frac{d\lambda_A}{d\theta} = 0 \Leftrightarrow f(\theta) = 0.$$

We now obtain:

$$\frac{df}{d\theta} = \pi \left[(1-\pi) \left(1 - 2 \cdot \frac{d\lambda_A}{d\theta} \right) - 2\pi \cdot \frac{d\varphi_A}{d\theta} \right],$$

and since $\frac{d\varphi_A}{d\theta} < 0$, this tells us that

$$\frac{d\lambda_A}{d\theta} \leq 0 \Rightarrow \frac{df}{d\theta} > 0 \quad \Rightarrow \quad f(\theta) \leq 0 \Rightarrow \frac{df}{d\theta} > 0.$$

But since f is continuous, the fact that it is increasing whenever it is negative and increasing when it crosses the zero line implies that it can only cross the zero line once. In other words, $f(\theta)$ can change signs at most once, going from negative to positive. But since $\frac{d\lambda_A}{d\theta}$ has the same sign, we conclude that λ_A must be convex: it decreases until some $\tilde{\theta}$, where $f(\tilde{\theta}) = 0$, and then increases. This, of course, provided that $\tilde{\theta} > 0$ — if not, then λ_A is strictly increasing.

We have concluded that λ_A is strictly increasing if, and only if, $f(0) \geq 0$. We now establish the conditions that ensure that. Solving $f(\theta) \geq 0$ gives us $(1-\pi)(\zeta - 2\pi\lambda_A) \geq 2\pi^2\varphi_A$, and using

(6), we can write this as

$$(1 - \pi)(\zeta - 2\pi\lambda_A) \geq \pi \left(\bar{w} - \frac{k}{\lambda_A} \right),$$

which yields the quadratic

$$2\lambda_A^2 - \left(\frac{\zeta}{\pi} - \frac{\bar{w}}{1 - \pi} \right) \lambda_A - \frac{k}{1 - \pi} \leq 0,$$

whose discriminant is

$$D = \left(\frac{\zeta}{\pi} - \frac{\bar{w}}{1 - \pi} \right)^2 + \frac{8k}{1 - \pi} > 0.$$

Since the smaller root is negative, the solution is at the larger root:

$$\widetilde{\lambda}_A = \frac{\frac{\zeta}{\pi} - \frac{\bar{w}}{1 - \pi} + \sqrt{D}}{4}.$$

The necessary and sufficient condition is that it is satisfied at $\theta = 0$, in which case:

$$\widetilde{\lambda}_A = \left(\frac{1}{4} \right) \left[\frac{\pi + c}{\pi} - \frac{\pi - c}{1 - \pi} + \sqrt{\left(\frac{\pi + c}{\pi} - \frac{\pi - c}{1 - \pi} \right)^2 + \frac{8k}{1 - \pi}} \right],$$

so the condition must obtain whenever

$$\lim_{\theta \rightarrow 0} \lambda_A \leq \widetilde{\lambda}_A$$

because the quadratic is a parabola and the solution is at the larger root. Thus, if this condition is satisfied, λ_A must be strictly increasing; otherwise, it will decrease first, and then increase.

We now show that $\lambda_A < \lambda_D$. First, we establish that λ_A is increasing as $\theta \Rightarrow \theta^*$. Observe that

λ_D is independent of θ , and recall that θ^* is such that (D) is satisfied with equality, which yields

$$\lim_{\theta \rightarrow \theta^*} f(\theta) = (1 - \pi)(\zeta - 2\pi\lambda_D) > 0,$$

because $\lambda_A \rightarrow \lambda_D$ and $\varphi_A \rightarrow 0$. Thus, λ_A is increasing when the anocratic equilibrium switches to the despotic one. Since λ_A is convex this implies that it can only possibly exceed λ_D as $\theta \rightarrow 0$. But this cannot be: the incentive to oppose when there is a positive probability of conflict is strictly weaker than when there is no such probability (even when retaliatory repression is at zero):

$$U_A(L; t) = \varphi_A W(t) + (1 - \varphi_A)(1 - t) - k < 1 - t - k = U_D(L; t),$$

where the inequality follows from the fact that any opponent must be some $t < 1/2 \Rightarrow t < 1 - t \Rightarrow W(t) = \pi(t - \theta) + (1 - \pi)(1 - t) < 1 - t$. If this type abstains, she would get $U_A(A; t) = \lambda_A(1 - t) + \lambda_A t$ in the anocratic equilibrium and $U_D(A; t) = \lambda_D(1 - t) + \lambda_D t$ in the despotic equilibrium. Thus, if $\lambda_A \geq \lambda_D$, the fact that $t < 1 - t$ would imply that $U_A(A; t) \geq U_D(A; t)$. Suppose now that $\lambda_A \geq \lambda_D$, which implies that $t_L(\lambda_A, \varphi_A) \geq t_L(\lambda_D, 0)$. Recall that $t_L(\lambda_A, \varphi_A)$ is the type that is precisely indifferent between opposing and abstaining, so

$$U_A(L; t_L(\lambda_A, \varphi_A)) = U_A(A; t_L(\lambda_A, \varphi_A)) \geq U_D(A; t_L(\lambda_A, \varphi_A)) \geq U_D(L; t_L(\lambda_A, \varphi_A)),$$

where the first inequality follows from the supposition that $\lambda_A \geq \lambda_D$ (per argument above), and the second inequality follows from the fact that $t_L(\lambda_D, 0)$ is the highest type to oppose in a despotic equilibrium, which implies that $t_L(\lambda_A, \varphi_A)$ cannot have a strict incentive to oppose. But this then

implies that $U_A(L; t_L(\lambda_A, \varphi_A)) \geq U_D(L; t_L(\lambda_A, \varphi_A))$, a contradiction to $U_A(L; t) < U_D(L; t)$.

Therefore, it must be that $\lambda_A < \lambda_D$ even as $\theta \rightarrow 0$, which establishes the claim. $\blacksquare$

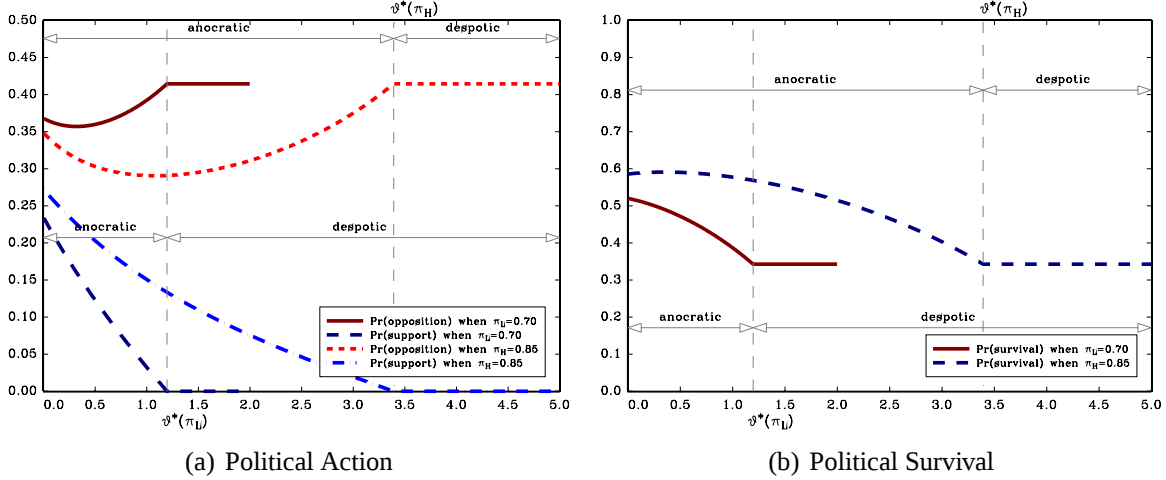


FIGURE 5: The Effect of Retaliatory Repression.

Parameters: $c = 0.1$, $k = 0.1$, and $\pi_L = 0.7$ (weak regime) or $\pi_H = 0.85$ (strong regime).

The relatively high value for π_L is necessary to ensure that Assumption 2 is satisfied despite θ being allowed to be relatively high.

Figure 5(a) illustrates the result from Lemma F for a weak and a strong regime. Note especially the fact that the probability of opposition in the anocratic equilibrium is always *lower* than the corresponding probability in the despotic equilibrium. It is easy to see that the latter must be constant because the retaliatory repression can only be imposed on the losing side when conflict occurs, and no conflict occurs in that equilibrium. In other words, retaliatory repression is essentially useless to a despot, and as result the probability of opposition is actually higher. This limits its usefulness as a policy tool. Consider the anocratic equilibrium where $\lambda_A < \lambda_D$ and $\varphi_A > 0$. Since the ruler's survival probability is decreasing in λ_A but increasing in φ_A , as evident from $\Omega_A = (1 - \lambda_A)^2 + 2\lambda_A\varphi_A \times \pi$, it follows that the ruler maximizes his chances of surviving by choosing some $\theta \in [0, \theta^*)$ and inducing the anocratic equilibrium. Figure 5(b) illustrates a case

where a strong regime chooses a strictly positive retaliatory repression level but the weak regime ends up with no penalties at all in the anocratic equilibrium.

Signaling Strength by Abandoning Repression

Consider a version of the model where the ruler knows the true probability of prevailing in a conflict, but the citizens do not. All other parameters, including any capacity constraints, are the same. Assume that the ruler can be either strong, in which case he wins the civil conflict with probability p_H , or weak, in which case he wins it with probability $p_L < p_H$. The citizens have a common belief $s \in (0, 1)$ that he is strong. If we let $\hat{s}$ denote the posterior belief after the ruler sets k , then the citizens' expected probability of him winning is $\pi = \hat{s} p_H + (1 - \hat{s}) p_L$. With this notation, Proposition 1, as well as lemmata 4 and 5, remain unchanged.

We now wish to ascertain whether it is possible to construct a separating equilibrium in which the ruler reveals his actual strength by choosing different levels of repression. To make the model interesting, assume that the capacity constraint, k_H , exceeds the despotic equivalent of k_L for the weak regime but not for the strong one. (For example, $k_H = k_H^2$ in Figure 3.) Consider now a strategy profile, in which the strong ruler induces the anocratic equilibrium by choosing the least-cost solution (k_L) and the weak one induces the despotic equilibrium by choosing at the capacity constraint (k_H). That is, (λ_A, φ_A) are the action probabilities when k_L is chosen and the citizens believe $\pi = p_H$, whereas $(\lambda_D, 0)$ are the action probabilities when k_H is chosen and citizens believe $\pi = p_L$.

It should be clear that the strong ruler has no incentive to change his strategy: he is getting the highest possible payoff in the anocratic equilibrium. The weak ruler, on the other hand, might be tempted to deviate because his expected payoff in the anocratic equilibrium where the citizens incorrectly attribute strength $\pi = p_H$ to him is strictly increasing. This is because these beliefs

induce supporters to turn out with a higher probability.

LEMMA G. *The weak ruler strictly benefits from citizens believing that he is strong.* $\square$

Proof of Lemma G. To see this, consider the probability of survival after this deviation from (20).

Taking the derivative with respect to π yields:

$$\frac{d\Omega_A(\pi; p_L)}{d\pi} = 2p_L\lambda_A \cdot \frac{d\varphi_A}{d\pi} - 2(1 - \lambda_A - p_L\varphi_A) \cdot \frac{d\lambda_A}{d\pi} > 0,$$

where we establish the inequality as follows. Using (14), we note that

$$2p_L\lambda_A \cdot \frac{d\varphi_A}{d\pi} = \left(\frac{p_L}{\pi}\right) \left[(\bar{w} - 2\pi\varphi_A) \cdot \frac{d\lambda_A}{d\pi} + (1 + \theta - 2\varphi_A)\lambda_A \right],$$

so we can rewrite the inequality above as

$$(1 + \theta - 2\varphi_A)\lambda_A > \left[2(1 - \lambda_A - p_L\varphi_A) - \left(\frac{p_L}{\pi}\right) (\bar{w} - 2\pi\varphi_A) \right] \cdot \frac{d\lambda_A}{d\pi}.$$

Since the proof of Lemma 8 establishes that $\frac{d\lambda_A}{d\pi} < 0$, it will be sufficient to show that the bracketed term is positive. Since $p_L < p_H = \pi$, it is sufficient to show that $2(1 - \lambda_A - p_L\varphi_A) > \bar{w} - 2\pi\varphi_A$, which can be written as $2 - 2\lambda_A - \bar{w} + 2(\pi - p_L)\varphi_A > 0$, which holds because $\lambda_A < 1/2$ and $\bar{w} < \pi < 1$. Thus, the weak ruler unequivocally benefits from the citizens believing he is strong. ■

The equilibrium can only be sustained if this temptation is not that alluring, as the following result shows.

PROPOSITION A. *Let k_L denote the lowest feasible cost of political action for both regimes, and let $k_H \in (\Delta(k_L; p_L), \Delta(k_L; p_H))$ denote their capacity constraint. The strategy profile in which the*

ruler chooses k_L when he is strong and k_H when he is weak is a separating equilibrium for any

$$p_L \leq \frac{(2 - \lambda_D - \lambda_A)(\lambda_A - \lambda_D)}{2\lambda_A\varphi_A}$$

irrespective of beliefs off the path of play. □

Proof of Proposition A. In an equilibrium, neither type wants to mimic the strategy of the other:

$$1 + \lambda_A^2 - 2(1 - p_H\varphi_A)\lambda_A \geq 1 + \lambda_D^2 - 2\lambda_D \quad (19)$$

$$1 + \lambda_D^2 - 2\lambda_D \geq 1 + \lambda_A^2 - 2(1 - p_L\varphi_A)\lambda_A. \quad (20)$$

Since $k_H < \Delta(k_L; p_H)$ by assumption, (19) holds with strict inequality, and the strong regime has no incentive to deviate. Rewriting (20) as specified in the proposition yields the condition that prevents the weak regime from deviating as well. The off-the-path beliefs are immaterial. The strong regime is at the highest possible survival probability in equilibrium already. If the weak regime deviates to any $k \in [k^*(p_H), \Delta(k_L; p_L))$, the payoff will be the same irrespective of the beliefs about π (because the despotic equilibrium prevails). If it deviates to any $k \in (k_L, k^*(p_H))$, then the most it can expect is that the citizens infer that the regime is strong, which would induce the anocratic equilibrium. But then choosing k_L maximizes the survival probability, so the only relevant deviation is to k_L , which the condition makes unprofitable. ■

The sufficient condition can be satisfied in two ways. First, one could fix p_H and make p_L small enough: in effect this ensures that however large the benefit from inducing the supporters to action under false pretenses, it will be outweighed by the fact that the ruler is actually unlikely

to prevail in the conflict their presence generates. For example, setting $p_L = 0.45$ and keeping the other parameters as in Figure 3 supports the separating equilibrium (for any $p_L \lesssim 0.48$). Second, one could fix p_L and reduce p_H enough: in effect this ensures that even if the ruler still has decent chances of prevailing in the conflict, the benefit from inducing the wrong beliefs is relatively small. For example, setting $p_H = 0.75$ and keeping the other parameters as in Figure 3 supports the separating equilibrium (any $p_L \lesssim 0.62$ will do).

It is worth noting that since we assumed repression to be costless to the ruler irrespective of regime strength, the separation is sustained by the riskiness of reducing repression: while the weak regime could exploit the benefit of supporters coming to its defense by feigning strength, it would still have to face its real, and not that great, odds of survival in the ensuing conflict. If it were the case that weak regimes also face higher costs of repression, then the incentive to permit separation would be diminished.

Differential Prevention

The original model assumes that preventive repression is equally costly for both anti-government and pro-government political actions. Suppose that the government did have some ability to discern whether the action is likely to be in its support and applied repression differently. To model this, assume that if preventive repression of level k is implemented, anti-government action still incurs a cost k , but pro-government action incurs a cost $\sigma k < k$, where $\sigma \in (0, 1]$ captures the government's ability to discriminate among political actions. Lower values of σ are associated with increasing ability, and so the original model, where the government is completely agnostic, is represented by $\sigma = 1$.

As before, $U_i(\text{Oppose}; t_i) > U_i(\text{Abstain}; t_i)$ whenever

$$2t_i < 1 - \frac{(\pi\theta + c)\varphi_{-i} + k}{1 - \lambda_{-i} - \pi\varphi_{-i}},$$

where we recall that since $\lambda_{-i} + \varphi_{-i} \leq 1$, it follows that $1 - \lambda_{-i} - \pi\varphi_{-i} > 0$. Thus, we recover the condition $t_i < t_L$ with

$$t_L(\lambda_{-i}, \varphi_{-i}) = \frac{1}{2} - \frac{(\pi\theta + c)\varphi_{-i} + k}{2(1 - \lambda_{-i} - \pi\varphi_{-i})}$$

from the original model.

Under the new assumption, $U_i(\text{Abstain}; t_i) > U_i(\text{Support}; t_i)$ whenever

$$2t_i < 1 + \frac{(1 - \pi)\theta + c}{\pi} + \frac{\sigma k}{\pi\lambda_{-i}},$$

which yields a slightly different version of $t_i < t'_R$ with

$$t'_R(\lambda_{-i}) = \frac{1}{2} + \left(\frac{1}{2\pi} \right) \left[(1 - \pi)\theta + c + \frac{\sigma k}{\lambda_{-i}} \right].$$

Since $t_L < 1/2 < t'_R$, it follows that the optimal strategies must be:

- $t_i < t_L$ plays $\lambda_i = 1$,
- $t_i \in (t_L, t'_R)$ plays $\lambda_i = \varphi_i = 0$, and
- $t_i > t'_R$ plays $\varphi_i = 1$.

Note that $t'_R < t_R$, so that the threshold for supporting the government is lower (it will be easier to satisfy) in the extended model.

Turning now to the next result, suppose that player $-i$ does not oppose the ruler in some equilibrium. This implies that player i would not support the ruler:

$$\lambda_{-i} = 0 \Rightarrow U(\text{Abstain}; t_i) = t_i > t_i - \sigma k = U(\text{Support}; t_i) \Rightarrow \varphi_i = 0.$$

But then $t_L(\lambda_i, 0) = 1/2 - k/(2(1 - \lambda_i))$, and since $\lambda_{-i} = 0$ requires that $t_L(\lambda_i, 0) \leq 0$, it follows that $\lambda_i \geq 1 - k > 0$ must obtain. In other words, it follows that player i must oppose the ruler with positive probability. In equilibrium, $\lambda_i = \Pr(t_i \leq t_L(0, \varphi_{-i}))$, so by the uniform distribution assumption it must be that $\lambda_i = t_L(0, \varphi_{-i})$, which implies that $t_L(0, \varphi_{-i}) \geq 1 - k$ must also hold. But this cannot be so because $k < 1$. Thus, we conclude that in any equilibrium, $t_L > 0$ for both players (that is, there exists no equilibrium in which some player does not oppose the government with positive probability). As before, $t_L > 0$ in every equilibrium.

Repression in the despotic form

Since the definition of t_L is the same as in the original model, the result for the despotic equilibrium goes through, and λ_D is the probability of an actor opposing the government. The existence threshold, however, is different. To see this, observe that we need to ensure $\lambda_i = 0$, or $t'_R(\lambda_{-i}) \geq 1$, which reduces to $\sigma k \geq \bar{w}\lambda_D$, or:

$$k \geq \bar{w}h'(\bar{w}) \equiv k^{*'} \in (0, 1),$$

where

$$h'(\bar{w}) = \frac{3\sigma + \bar{w} - \sqrt{(3\sigma + \bar{w})^2 - 8\sigma^2}}{4\sigma^2}.$$

Some algebra further shows that

$$\frac{dk^{*'}}{d\sigma} < 0,$$

and that

$$\lim_{\sigma \rightarrow 1} k^{*'} = k^* \quad (\text{by inspection})$$

$$\lim_{\sigma \rightarrow 0} k^{*'} = 1 \quad (\text{repeated application of L'Hôpital's rule})$$

These imply that $k^{*'} > k^*$ for all $\sigma < 1$. The threshold for the despotic form under selective repression is always *greater* than the threshold under indiscriminate repression — and so the range of values of k for which the equilibrium takes that form is *smaller* — and this threshold is *increasing* as the government becomes better able to discriminate (σ goes down). In other words, as the

government's ability to target preventive repression against potential opponents improves, it can mobilize its supporters better, which in turn means that the equilibrium will take the anocratic form at levels of repression that previously resulted in the despotic form.

Repression in the anocratic form

Turning now to the anocratic form, we can write the system of equations as

$$3\lambda - 2\lambda^2 - 2\pi\lambda\varphi = 1 - k - \zeta\varphi \quad (21)$$

$$2\pi\lambda\varphi = \bar{w}\lambda - \sigma k, \quad (22)$$

where $\zeta = (1 + \theta)\pi + c > \pi$ as before. This yields the cubic:

$$G'(\lambda) = -2\lambda^3 + (3 - \bar{w})\lambda^2 - \left(1 - (1 + \sigma)k - \frac{\bar{w}\zeta}{2\pi}\right)\lambda - \frac{\zeta\sigma k}{2\pi} = 0. \quad (23)$$

Some algebra analogous to what we used in the analysis of the original model shows that this cubic has a unique solution, $\lambda'_A \in (0, 1/2)$ if, and only if, $k < k^*$. This, in turn, yields the unique value for $\varphi'_A \in (0, 1/2)$ as well. Thus, the original result is recovered: the game has a unique equilibrium that takes the anocratic form if, and only if, $k \leq k^*$, and the despotic form otherwise.

We now show that λ'_A is monotonic. Recall that at the optimum, $G'(\lambda'_A) = 0$. Using (23), this yields

$$-2\lambda'^2_A + (3 - \bar{w})\lambda'_A - \left[1 - (1 + \sigma)k - \frac{\bar{w}\zeta}{2\pi}\right] = \frac{\zeta\sigma k}{2\pi\lambda'_A}. \quad (24)$$

Recall that at the optimum,

$$\left. \frac{dG'}{dk} \right|_{\lambda=\lambda'_A} = \left. \frac{\partial G'}{\partial \lambda} \right|_{\lambda=\lambda'_A} \cdot \left. \frac{d\lambda}{dk} \right|_{\lambda=\lambda'_A} + \left. \frac{\partial G'}{\partial k} \right|_{\lambda=\lambda'_A} = 0.$$

Since

$$\left. \frac{\partial G'}{\partial \lambda} \right|_{\lambda=\lambda'_A} = -6\lambda'^2_A + 2(3 - \bar{w})\lambda'_A - \left[1 - (1 + \sigma)k - \frac{\bar{w}\zeta}{2\pi} \right]$$

and, using (24),

$$= (3 - \bar{w} - 4\lambda'_A)\lambda'_A + \frac{\zeta\sigma k}{2\pi\lambda'_A} > 0,$$

it follows that,

$$\operatorname{sgn} \left(\frac{d\lambda'_A}{dk} \right) = -\operatorname{sgn} \left(\frac{\partial G'}{\partial k} \right).$$

Let $f(k) = \frac{\partial G'}{\partial k}$, so that

$$f(k) = (1 + \sigma)\lambda'_A - \frac{\zeta\sigma}{2\pi}.$$

Since

$$\frac{df}{dk} = (1 + \sigma) \cdot \frac{d\lambda'_A}{dk},$$

we obtain:

$$\operatorname{sgn} \left(\frac{df}{dk} \right) = \operatorname{sgn} \left(\frac{d\lambda'_A}{dk} \right) = -\operatorname{sgn}(f(k)).$$

In other words if $f(k)$ is increasing, it must be that $f(k) < 0$, and if $f(k)$ is decreasing it must

be that $f(k) > 0$. This implies that $f(k)$ must be monotonic. To see this, suppose that $f(k)$ is increasing for some k , and thus $f(k) < 0$ must hold. Now increase k and note that as long as $f(k) < 0$, the function must monotonically keep increasing until it gets to 0. But now if it were to keep increasing, $f(k) > 0$, a contradiction because this implies that it should be decreasing in k . If, on the other hand, it were to decrease, then $f(k) < 0$, also a contradiction because this means it should be increasing in k . A similar exercise starting with $f(k) > 0$, and thus decreasing shows that the function cannot switch signs (or reach zero).

Since λ'_A is monotonic and $\lambda'_A(k^{*'}) = \lambda_D(k^{*'})$, the sign of $f(k^{*'})$ will tell us whether λ'_A is increasing or decreasing, so

$$f(k^{*'}) = (1 + \sigma)\lambda_D(k^{*'}) - \frac{\zeta\sigma}{2\pi}.$$

This yields the analogue to condition (P), which we can write as:

$$\sigma \left[2 \left(\theta + \frac{c}{\pi} \right) - 1 \right] + (1 + \sigma) \sqrt{1 + 8\bar{w}h'(\bar{w})} > 3. \quad (P')$$

We conclude that if condition (P') is satisfied, then λ'_A is increasing in k , otherwise it is decreasing. When it comes to the equilibrium probabilities of support and opposition, we have recovered all results from the original model.

We now show that, as one would expect, increasing the government's ability to discriminate with preventive repression has a positive effect on the probability that its supporters become active. Since both (21) and (22) must hold in equilibrium, we differentiate both sides with respect to σ to

obtain:

$$\left(3 - 4\lambda'_A - 2\pi\varphi'_A\right) \cdot \frac{d\lambda'_A}{d\sigma} = -\left(\zeta - 2\pi\lambda'_A\right) \cdot \frac{d\varphi'_A}{d\sigma} \quad (25)$$

$$-\left(\bar{w} - 2\pi\varphi'_A\right) \cdot \frac{d\lambda'_A}{d\sigma} + k = -2\pi\lambda'_A \cdot \frac{d\varphi'_A}{d\sigma} \quad (26)$$

Since $3 - 4\lambda'_A - 2\pi\varphi'_A > 0$ and $\zeta - 2\pi\lambda'_A > 0$, (25) implies that

$$\frac{d\lambda'_A}{d\sigma} \geq 0 \Rightarrow \frac{d\varphi'_A}{d\sigma} < 0.$$

Since $\bar{w} - 2\pi\varphi'_A > 0$, (26) further implies that

$$\frac{d\lambda'_A}{d\sigma} \leq 0 \Rightarrow \frac{d\varphi'_A}{d\sigma} < 0,$$

we conclude that

$$\frac{d\varphi'_A}{d\sigma} < 0. \quad (27)$$

Since (25) tells us that

$$\text{sgn}\left(\frac{d\lambda'_A}{d\sigma}\right) = -\text{sgn}\left(\frac{d\varphi'_A}{d\sigma}\right),$$

equation (27) further implies that

$$\frac{d\lambda'_A}{d\sigma} > 0. \quad (28)$$

In other words, as σ decreases (so the government's repression targets its potential opponents without hurting its potential supporters), the probability that the supporters act on its behalf increases, whereas the probability that its opponents become active decreases.

Survival probability in the anocratic form

Consider now the effect of repression in on the survival probability in the extended model:

$$\frac{d\Omega'_A}{dk} = 2 \left[\pi\lambda'_A \cdot \frac{d\varphi'_A}{dk} - (1 - \lambda'_A - \pi\varphi'_A) \cdot \frac{d\lambda'_A}{dk} \right],$$

and thus:

$$\text{sgn} \left(\frac{d\Omega'_A}{dk} \right) = \text{sgn} \left(\pi\lambda'_A \cdot \frac{d\varphi'_A}{dk} - (1 - \lambda'_A - \pi\varphi'_A) \cdot \frac{d\lambda'_A}{dk} \right). \quad (29)$$

Using (21) and (22), define

$$(3 - 4\lambda'_A - 2\pi\varphi'_A) \cdot \frac{d\lambda'_A}{dk} + 1 = -(\zeta - 2\pi\lambda'_A) \cdot \frac{d\varphi'_A}{dk} \quad (30)$$

$$-(\bar{w} - 2\pi\varphi'_A) \cdot \frac{d\lambda'_A}{dk} + \sigma = -2\pi\lambda'_A \cdot \frac{d\varphi'_A}{dk}. \quad (31)$$

As before, these imply that

$$\frac{d\varphi'_A}{dk} < 0,$$

which in turn means that

$$\frac{d\lambda'_A}{dk} \geq 0 \Rightarrow \frac{d\Omega'_A}{dk} < 0.$$

That is, if condition (P') is satisfied (so λ'_A is increasing in k), then the survival probability must be strictly decreasing in the anocratic equilibrium, just as it is in the original model.

Consider them the case when condition (P') fails, so λ'_A is strictly decreasing. We differentiate

both sides of equations (30) and (31) again to obtain

$$4 \left(\frac{d\lambda'_A}{dk} \right)^2 + 4\pi \cdot \frac{d\lambda'_A}{dk} \cdot \frac{d\varphi'_A}{dk} - (3 - 4\lambda'_A - 2\pi\varphi'_A) \cdot \frac{d^2\lambda'_A}{dk^2} = (\zeta - 2\pi\lambda'_A) \cdot \frac{d^2\varphi'_A}{dk^2} \quad (32)$$

$$(\bar{w} - 2\pi\varphi'_A) \cdot \frac{d^2\lambda'_A}{dk^2} - 4\pi \cdot \frac{d\lambda'_A}{dk} \cdot \frac{d\varphi'_A}{dk} = 2\pi\lambda'_A \cdot \frac{d^2\varphi'_A}{dk^2}. \quad (33)$$

We now show that if λ'_A is decreasing, then it must be convex. Assume that $\frac{d\lambda'_A}{dk} < 0$. If $\frac{d^2\varphi'_A}{dk^2} \leq 0$, then the right-hand side of (32) is non-positive, and since the first and second terms on the left-hand side are strictly positive, the only way the equality can obtain is if $\frac{d^2\lambda'_A}{dk^2} > 0$. In other words,

$$\frac{d^2\varphi'_A}{dk^2} \leq 0 \Rightarrow \frac{d^2\lambda'_A}{dk^2} > 0.$$

If $\frac{d^2\varphi'_A}{dk^2} \geq 0$, then the right-hand side of (33) is non-negative, and since the second term on the left-hand side is strictly negative, the only way the equality can obtain is if $\frac{d^2\lambda'_A}{dk^2} > 0$. In other words,

$$\frac{d^2\varphi'_A}{dk^2} \geq 0 \Rightarrow \frac{d^2\lambda'_A}{dk^2} > 0.$$

Thus, we conclude that λ'_A is decreasing at decreasing rates:

$$\frac{d\lambda'_A}{dk} < 0 \Rightarrow \frac{d^2\lambda'_A}{dk^2} > 0.$$

Using (31) we can write

$$\frac{d\Omega'_A}{dk} = 2\pi\lambda'_A \cdot \frac{d\varphi'_A}{dk} - 2(1 - \lambda'_A - \pi\varphi'_A) \cdot \frac{d\lambda'_A}{dk}$$

$$\begin{aligned}
&= (\bar{w} - 2\pi\varphi'_A) \cdot \frac{d\lambda'_A}{dk} - \sigma - 2(1 - \lambda'_A - \pi\varphi'_A) \cdot \frac{d\lambda'_A}{dk} \\
&= -[2(1 - \lambda'_A) - \bar{w}] \cdot \frac{d\lambda'_A}{dk} - \sigma,
\end{aligned}$$

where we note that

$$2(1 - \lambda'_A) - \bar{w} > 1 - \bar{w} > 0,$$

where the first inequality follows from $\lambda'_A < 1/2$, and the second from $\bar{w} = \pi - c < 1$.²⁸ From this, it follows that

$$\frac{d^2\Omega'_A}{dk^2} = 2 \left(\frac{d\lambda'_A}{dk} \right)^2 - [2(1 - \lambda'_A) - \bar{w}] \cdot \frac{d^2\lambda'_A}{dk^2}.$$

We now wish to show that the second derivative is negative, which we can express as follows (after multiplying both sides by 2):

$$(4 - 4\lambda'_A - 2\bar{w}) \cdot \frac{d^2\lambda'_A}{dk^2} - 4 \left(\frac{d\lambda'_A}{dk} \right)^2 > 0.$$

Adding (32) and (33) yields:

$$(3 - 4\lambda'_A - \bar{w}) \cdot \frac{d^2\lambda'_A}{dk^2} - 4 \left(\frac{d\lambda'_A}{dk} \right)^2 = -\zeta \cdot \frac{d^2\varphi'_A}{dk^2}.$$

Since $4 - 4\lambda'_A - 2\bar{w} - (3 - 4\lambda'_A - \bar{w}) = 1 - \bar{w} > 0$, it follows that the desired inequality must obtain whenever $\frac{d^2\varphi'_A}{dk^2} \leq 0$.

28. It immediately follows that

$$\frac{d\lambda'_A}{dk} \geq 0 \Rightarrow \frac{d\Omega'_A}{dk} < 0,$$

which we have already established.

Making the appropriate substitutions in (32) and (33) and simplifying yields:

$$\begin{aligned} & \left[2\pi\lambda'_A + \frac{\zeta - 2\pi\lambda'_A}{3 - 4\lambda'_A - 2\pi\varphi'_A} \right] \cdot \frac{d^2\varphi'_A}{dk^2} \\ &= 4 \left(\frac{\bar{w} - 2\pi\varphi'_A}{3 - 4\lambda'_A - 2\pi\varphi'_A} \right) \left[\frac{d\lambda'_A}{dk} + \pi \left(1 - \frac{3 - 4\lambda'_A - 2\pi\varphi'_A}{\bar{w} - 2\pi\varphi'_A} \right) \cdot \frac{d\varphi'_A}{dk} \right] \cdot \frac{d\lambda'_A}{dk} \end{aligned}$$

The bracketed term on the left-hand side is strictly positive, as are the first two terms on the right-hand side. Since $\frac{d\lambda'_A}{dk} < 0$, it follows that if the bracketed term on the right-hand side is positive, the right-hand side must be negative, which would imply that $\frac{d^2\varphi'_A}{dk^2} < 0$ as required. Thus, we need to show that

$$\frac{d\lambda'_A}{dk} + \pi \left(1 - \frac{3 - 4\lambda'_A - 2\pi\varphi'_A}{\bar{w} - 2\pi\varphi'_A} \right) \cdot \frac{d\varphi'_A}{dk} > 0,$$

or, after a bit of algebra, that

$$-\pi (3 - 4\lambda'_A - \bar{w}) \cdot \frac{d\varphi'_A}{dk} > -(\bar{w} - 2\pi\varphi'_A) \cdot \frac{d\lambda'_A}{dk},$$

which, after using (31), can be written as

$$\pi (3 - 6\lambda'_A - \bar{w}) \cdot \frac{d\varphi'_A}{dk} < \sigma. \quad (34)$$

Observe now that

$$\lambda'_A(k^{*'}) = \lambda'_D = \frac{3 - \sqrt{1 + 8k^{*'}}}{4} < \frac{3 - \bar{w}}{6},$$

which means that $3 - 6\lambda'_A(k^{*'}) - \bar{w} > 0$, and thus (34) is satisfied at $k^{*'}$ (because the left-hand side is negative). Since the derivative is monotonic, this establishes the sign everywhere.

We conclude that Ω_A is concave in k , which implies that there exists $\hat{k} \in [0, k^{*'}]$ that maximizes it. Since the probability of survival is strictly increasing for $k < \hat{k}$ and strictly decreasing for $k > \hat{k}$, we now establish conditions that ensure that it will be monotonic over the admissible range.

Assume now an interior optimum at $\hat{k} \in (0, k^{*'})$, where

$$\frac{d\Omega'_A}{dk} = \frac{\partial\Omega'_A}{\partial\lambda'_A} \cdot \frac{d\lambda'_A}{dk} + \frac{\partial\Omega'_A}{\partial\varphi'_A} \cdot \frac{d\varphi'_A}{dk} = 0,$$

let $\Omega_A^* \equiv \Omega'_A(\hat{k})$, and consider how changing the ability to discriminate will alter it:

$$\begin{aligned} \frac{d\Omega_A^*}{d\sigma} &= \left[\frac{\partial\Omega'_A}{\partial\lambda'_A} \cdot \frac{d\lambda'_A}{dk} + \frac{\partial\Omega'_A}{\partial\varphi'_A} \cdot \frac{d\varphi'_A}{dk} \right] \cdot \frac{d\hat{k}}{d\sigma} + \frac{\partial\Omega'_A}{\partial\lambda'_A} \cdot \frac{d\lambda'_A}{d\sigma} + \frac{\partial\Omega'_A}{\partial\varphi'_A} \cdot \frac{d\varphi'_A}{d\sigma} \\ &= \frac{\partial\Omega'_A}{\partial\lambda'_A} \cdot \frac{d\lambda'_A}{d\sigma} + \frac{\partial\Omega'_A}{\partial\varphi'_A} \cdot \frac{d\varphi'_A}{d\sigma} < 0, \end{aligned}$$

where the inequality follows from

$$\frac{\partial\Omega'_A}{\partial\lambda'_A} = -2(1 - \lambda'_A - \pi\varphi'_A) < 0, \quad \frac{\partial\Omega'_A}{\partial\varphi'_A} = 2\pi\varphi'_A > 0,$$

and (27) and (28), which tell us that

$$\frac{d\lambda'_A}{d\sigma} > 0, \quad \frac{d\varphi'_A}{d\sigma} < 0.$$

In other words, the optimal survival probability decreases as the ability to discriminate gets worse.

Recall that at an interior solution, $\hat{k}$ solves

$$-(2 - \bar{w} - 2\lambda'_A(\hat{k})) \cdot \frac{d\lambda'_A}{dk} \Big|_{k=\hat{k}} = \sigma.$$

Taking the derivative of both sides with respect to σ yields

$$\left[2 \left(\frac{d\lambda'_A}{dk} \right)^2 - [2(1 - \lambda'_A) - \bar{w}] \cdot \frac{d^2\lambda'_A}{dk^2} \right] \cdot \frac{d\hat{k}}{d\sigma} = 1,$$

or

$$\frac{d^2\Omega'_A}{dk^2} \cdot \frac{d\hat{k}}{d\sigma} = 1.$$

But then the concavity of Ω'_A tells us that repression must go down:

$$\frac{d^2\Omega'_A}{dk^2} < 0 \Rightarrow \frac{d\hat{k}}{d\sigma} < 0.$$

In other words, as the ability to discriminate gets worse, the optimal (interior) repression decreases.

We know from the original model that $\sigma = 1$ implies that $\hat{k} \rightarrow 0$. Thus, the above result implies that there exists $\bar{\sigma} < 1$ such that for all $\sigma > \bar{\sigma}$, the ruler must choose the lowest feasible level of repression because the survival probability is strictly decreasing for all positive values of repression. Moreover, there exists $\underline{\sigma} > 0$ such that $\hat{k} = k'$ for all $\sigma < \underline{\sigma}$. In other words, as σ becomes sufficiently small, the ruler will maintain the maximum feasible repression.

Recall now that $\Omega'_A = (1 - \lambda'_A)^2 + 2\lambda'_A\varphi'_A \times \pi$, which means that

$$\frac{d\Omega'_A}{d\sigma} = 2 \left[\pi\lambda'_A \cdot \frac{d\varphi'_A}{d\sigma} - (1 - \lambda'_A - \pi\varphi'_A) \cdot \frac{d\lambda'_A}{d\sigma} \right] < 0,$$

and thus, predictably, the ruler's chances of survival get better (Ω_A is higher) as his ability to discriminate with repression improves (σ is lower).

The survival probability in the despotic equilibrium is the same in the extended model as in the original (because the supporters are not active, and the opponents pay the full cost of preventive repression): $\Omega'_D = \Omega_D$. Thus, for $k \geq k^{*'}$, the payoff for the ruler is the same in both cases. For $k \in (k^*, k^{*'})$, the equilibrium would still be despotic in the original model but would take the anocratic form with discriminatory repression. In this range, the ruler's payoff is still increasing in k in the original model but decreasing in the extended one. For $k \leq k^*$, the equilibrium takes the anocratic form in both models, and the ruler's payoff is decreasing in k . The ruler's payoff for $k < k^{*'}$ in the extended model is strictly greater than his payoff in the original model.

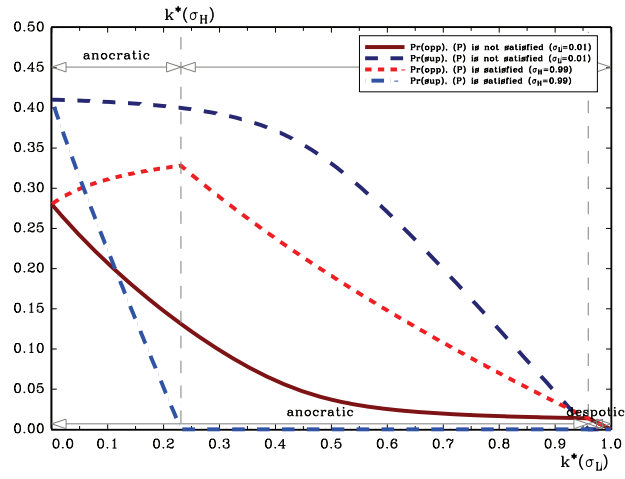
We conclude that the authoritarian wager exists in the extended model, and that it is, in fact, even more likely to occur because the ruler's payoff in the anocratic equilibrium is strictly higher but in the despotic the same, it follows that the despotic equivalent to no repression must be strictly greater in the extended model as well. In other words, regimes that have better abilities to limit the negative effects of preventive repression on their supporters are, in fact, more likely to take the authoritarian wager because they can rely on even marginal supporters to not be deterred in acting on their behalf. These regimes are much more likely to come out ahead with the wager as well.

Figures 6, 7, and 4 show the effects of repression for various levels of discriminatory capacity. Figure 6 plots the probabilities of political action and survival for the opposite cases of near perfect capacity (where the chance of incorrectly repressing a supporter is merely $\sigma = 0.01$), and incapacity close to the original model (where this chance is $\sigma = 0.99$). Observe that for the given parameter configuration (P') is not satisfied in the high capacity case but is satisfied in the low

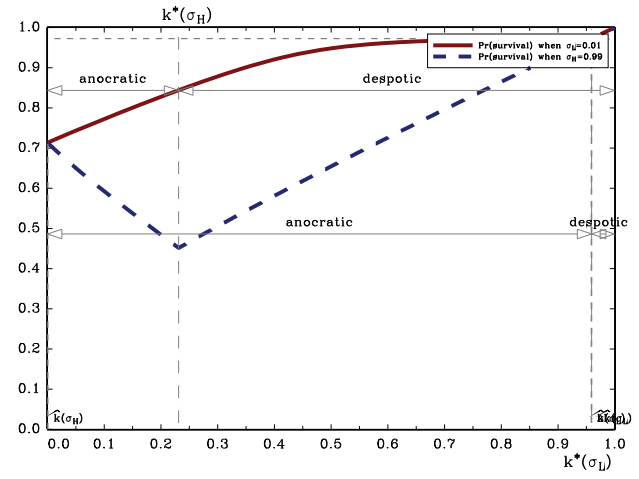
capacity case. The anocratic region extends almost over the entire range of repression when the government has high capacity to differentiate, and the ruler will always pick the highest possible preventive repression.

Figure 7 compares somewhat more realistic capacities, where the chance of incorrectly repressing a supporter are 25% and 85%, respectively. When the government is still relatively limited in its ability to differentiate *ex ante*, the original result is fully recovered, and a reduction in repressive capacity results in collapse of repression. The authoritarian wager is not as stark when the government has fairly good differentiation capacity but it still exists. For instance, if the repression capacity is reduced from, let's say, $k_H^1 = 0.7$ (where $k = k_H^1$ in the despotic equilibrium) to $k_H^2 = 0.55$, then preventive repression will fall to the interior optimum in the anocratic equilibrium (about $k = 0.15$). The wager is attenuated but clearly still there.

Finally, Figure 4 (in the paper) compares that fairly good capacity, $\sigma = 0.25$, with one that is quite high, $\sigma = 0.10$. The interior optimum in the anocratic equilibrium goes up, as established in the proofs, when the government becomes better able to differentiate. Correspondingly, if repression capacity is reduced from, say, $k_H^1 = 0.80$ (where $k = k_H^1$ in a despotic equilibrium) to $k_H^2 = 0.65$, then ruler will respond with a weak version of the wager at the interior optimum (about $k = 0.55$). If the capacity falls below that optimum, then the wager will cease to exist: the ruler will simply repress at the maximum capacity (the equilibrium will still take the anocratic form). Observe, however, just how high the differentiation capacity has to be and how drastic a fall in repression capacity must occur before the wager is completely eliminated.



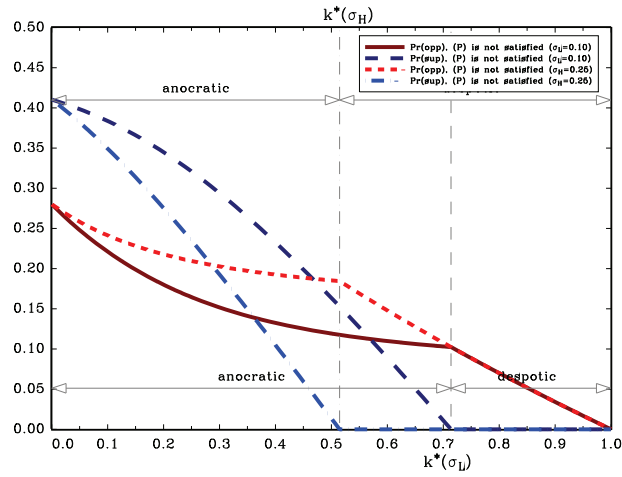
(a) Political Action



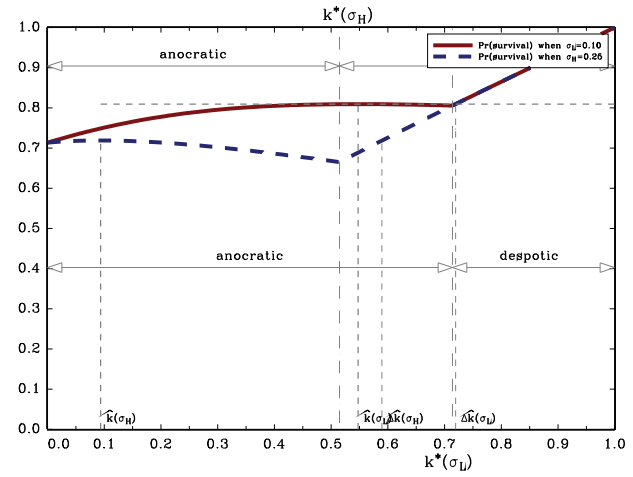
(b) Political Survival

FIGURE 6: The Effect of Discriminatory Capacity.

Parameters: $c = 0.1$, $\theta = 0.35$, and $\pi = 0.85$.



(a) Political Action



(b) Political Survival

FIGURE 7: The Effect of Discriminatory Capacity.

Parameters: $c = 0.1$, $\theta = 0.35$, and $\pi = 0.85$.