

Federalisms: Unstable by Design

Jenna Bednar*
Olin Fellow, U.S.C. Law School
jbednar@law.usc.edu

April 15, 1997

Abstract

Federal systems are crippled by power grabbing between central and regional governments, as well as burden-shifting schemes between regions. Existing models of federalisms assume regional diversity to account for inter-regional tension. However, these models set aside entirely the problem of inter-level competition.

This paper presents a unified framework for understanding threats to federal stability. The model's $n + 1$ structure accomodates both dimensions of federal instability. Furthermore, this paper is able to offer a theoretical alternative to explanations of instability that rely upon regional diversity or citizen patriotism; identically selfish preferences, in the decentralized setting, can generate instability. Additionally, under certain institutional conditions, the paper offers an equilibrium that embraces the persistence of competition in a stable federation.

1 Introduction

That at least some federal systems are unstable comes as no surprise to any reader of the newspaper. Canada, Russia, Bosnia, Israel, and Northern Ireland remind us that intergovernmental rivalries sicken federal and

*Many minds have influenced this work: I thank particularly John Ferejohn, Brian Gaines, Scott Page, Jack Rakove, Matt Spitzer, and Barry Weingast, and seminar participants at Stanford, Caltech, UC-Irvine, SUNY Stony Brook, the University of Pittsburgh, Iowa, and the Claremont Graduate School. Some very good ideas emerged in these discussions; doubtless, had I heeded all of your suggestions, the paper would be improved. I thank the Olin Foundation for research support.

quasi-federal¹ systems, at best making their future uncertain, and at worst, erupting into violent civil war. To date, the cause of federal instability has not been nailed down in a general sense: we treat failing federalisms on a case-by-case basis, diagnosing problems as they arise. Prescriptions for stability follow logically from this ad hoc approach, treating symptoms, rather than the disease.

Inter-regional rivalry is not the only cause of federal distress, and diversity is misleading; in fact, heterogeneous preferences could be a symptom of a deeper problem, aggravated by another, more general cause. This paper deviates from the traditional study of federal stability, to consider the effect of federalism's decentralized structure. The paper will argue that all federalisms are susceptible to destabilizing intergovernmental competition, regardless of the alignment of preferences.

Regional diversity proves to be a popular and compelling explanation for instability.² We know, however, that federalism is prescribed precisely when the regions represent clumpings of heterogeneous preferences.³ Therefore, an excuse for instability that relies upon heterogeneity begs the question: if we implement a federalism to satisfy regional differences, then how do we know when the diversity passes from necessitating federation, to rupturing it? In other words, why does the ethnic rivalry between the Muslim and the Croats render Bosnian federation untenable, while the differences between French and English Canadians, or Czech and Slovak Europeans, can be peaceably settled, and French and German Swiss manage to govern themselves harmoniously? Diversity cannot be the litmus test for federal failure, since it also seems to be the reason for establishment of the federation, and it is present in thriving federations. At the least, a heterogeneity of interests is not a sufficient condition for instability, and it might not be necessary.

Others have proposed that citizen loyalty can be the glue that holds a federal union together. William Riker⁴ suggests that a transference of loyalty from region to union might be the most important element to a stable federation. Presumably, if all citizens want the union to work, they will not pursue actions that jeopardize it. When citizens do not believe in the federation, the federation will fail. Riker's argument lies in good

¹A quasi-federal system devolves some authority, asymmetrically, to targeted regions.

²See, for example, Franck 1968, Friedrich 1968, Hicks 1978, Lemco 1991.

³See, for example, Tiebout (1956), Oates (1972), Ostrom (1991), Peterson (1995), Tullock (1969); see Cremer & Palfrey 1996, 1997 for a mathematical exercise of this maxim.

⁴1964, especially p. 111.

company,⁵ but the argument's reasonableness dissipates when we press for a micromotive foundation: we don't know what initiates the commitment to the union, nor do we have an idea of what causes the interest in the union to disintegrate.

James Madison points toward the weakness in the loyalty argument: in the course of musing on the incapacity of the Articles of Confederation to provide for stable united governance, he wondered how anyone could have believed that "a treaty of amity of commerce and of alliance" could have mustered the strength and unity of purpose required for longevity. Madison reckons responsible a zealous optimism; leadership blinded by lusty hope to the natural tendencies of men and their governments to compete with one another. He cites:

a mistaken confidence that the justice, the good faith, the honor, the sound policy, of the several legislative assemblies would render superfluous any appeal to the ordinary motives by which the laws secure the obedience of individuals: a confidence which does honor to the enthusiastic virtue of the compilers, as much as the inexperience of the crisis apologizes for their errors. (Madison, "The Vices of the Political System of the United States," in Rutland et al, eds., *The Papers of James Madison*, vol. 9, 1975).

Madison quashes our hope that honest intentions will save a federal union, labelling such arguments naive and gullible.

This essay turns to a rival explanation: The decentralized decision-making structure itself contributes to the instability of a federation. In other words, a federalism is unstable by design. This paper moves us away from explanations of instability that rely upon diverse preferences between regions, and also argues that a commitment to the union does not make it work. In this paper, I will argue that the decentralized decision-making process introduces a fault to the federal system; instability is inherent.

Three sections remain to this paper. I first present the signs of instability in a federal system, and discuss the current state of the political economy literature as it relates to these symptoms. Section 3 presents the model and the main results, and provides the intuition for the application of the model to federal stability. In Section 4 I conclude.

⁵See Elazar (1987), Franck (1968) and Beer (1995).

2 The Symptoms of Instability

The specific problems of a federal union, as opposed to a unitary system, involve intergovernmental competition between the separate regions, and between layers of government.⁶ Regions argue between themselves and through the federal institutional channels, and regions and the central government struggle for authority. All federalisms are familiar with these symptoms.

Membership in a union requires sacrifices and compromises: regions (the sub-level of government, such as states or provinces) try to shift these burdens of federal membership onto the backs of other regions, by bickering over tax incidence, trade arrangements, location of industry, redistribution, and the like. Perhaps guided by Madison, who suspected that the tendency to federal instability was “rather to anarchy among the members, than to tyranny in the head,” (Fed. 18) a healthy portion of the federalism literature is consumed with interregional rivalry of this sort. The fiscal federalism literature supports legions of economists who crank out Tiebout model after Tiebout model to generate optimal taxation and distribution schemes. Legal scholars and political scientists also warn of the problems caused by regional competition and burden-shifting, and related syndromes.⁷ Political economists tend to hone in on competition between regions.⁸ These theories can be generally described as n -agent models (for the n regions); commonly, dissent is driven by an assumption of heterogeneity of preferences between the regions.

The political economy literature tends to overlook the contribution of the central government to instability. The central government is not an angelic constraint;⁹ it too is guilty of contributing to the risk associated with federal association. Some of political burden-shifting battles seep into the central government’s agenda: charges of favoritism and bias disturb a federal balance as much as battles between regions. With some tweaks to the

⁶Since we are interested in instability in federal systems, we can exclude all destabilizing factors that can affect any polity regardless of constitutional design. In general, external or internal military threat and economic and civil strife fall into this category. However, one point we will want to consider is that weak federalisms might be more susceptible.

⁷For example, see Madison’s Notes on the Vices and Hamilton in Federalist Nos. 6-8. In more recent work, see for example Enrich (1996) and Brown (1995) as well as the cites above.

⁸Maggi unpub; Persson & Tabellini 1996a, 1996b; Cremer & Palfrey 1996, 1997; Caplan 1996; Weingast 1994a, 1994b, 1995.

⁹For indications of its usefulness, however, see Wechsler 1954, Weingast 1995, Montinola, Qian, and Weingast 1995.

existing n -agent models to allow for capture at the central level by regional agents, the standard model of a federalism, based upon diverse preferences between regions, can be maintained.¹⁰

A more complete model of federal instability would include the power tug-of-war between central and regional governments. Centralization or peripheralization, when carried to extremes, threatens federalisms as much as burden-shifting,¹¹ but the standard n -agent models of federal relations cannot deal effectively with its causes: encroachment and shirking. Both implicate the federal government directly in destabilizing action and suggest that the federal government is a strategic player, and should be modelled as such.

Burden-shifting, encroachment, favoritism, and shirking are all actions taken by one governmental agent that affect the other agents' value of the union. Since agents cannot effectively monitor one another's behavior, federalism suffers from a moral hazard problem. The cause of the moral hazard, I will argue, comes directly from the decentralized design of the federalism, and not necessarily from any problems of diversity or dishonesty. It is common to all federal systems.

3 Constructing a 'Pure' Model of a Federalism

Federal structures are first and foremost decentralized decision-making systems. Governmental authority is parsed out between levels of government.¹² Following the definition provided by Riker (1964), at a minimum, both regional and central governments enjoy complete authority in at least one jurisdiction.

¹⁰GET AFRICA CITE and Dixit & Londregan 1996.

¹¹See, for example, Riker's (1964) judgment that over-peripheralized federations cannot survive, and Hamilton's similar warning, with the words *divide et impera* to convince by fear the unconvinced of federalism's virtues. Bednar et al (1995) argue that the mere threat of centralization can destabilize a union: see the section on Canada.

¹²For simplicity, I restrict the problem to two levels. The model can be extended to include three or more levels by breaking the federalism into bi-level games. For example, in a federalism with a central, regional, and local level, the regional agents might be motivated to protect local interests in games with the central level (as they are in this essay), while in interactions with the local level the regional agents act as unifiers. A potential complication arises from a central-local game, which skips the intermediate level altogether. While I do not address this problem in this essay, it merits further consideration, as we see examples of it in regional development projects, such as Europe's structural funds, to develop infrastructure in designated regions.

I impose an additional condition, which Riker might have implied but did not express explicitly: the center and regions must be electorally independent. Agents at either regional or central level obtain office through constituent election, not through appointment from one to the other. In this model, I assume that citizens are divided equally and exhaustively into mutually exclusive regions. Each region elects an agent, and the collective citizenry elect one central agent, generating $n + 1$ total agents in the union.

I make no other assumptions about the structure of a federalism. Ordinarily, onto this decentralized skeleton constitutional framers will hang all sorts of institutional trappings, such as bicameralism, separation of powers, and an independent court. Other institutions, like a party system, tend to develop with time and impact the relations between the levels. But these modifications vary across federations, and since I am trying to construct a general hypothesis of the source of internal conflict, I will set these modifications aside for the moment.¹³ For the moment, consider an “institution-free” federalism, where no corrective institutions are available.

A federation combines the strengths of coordination with the local satisfaction of limited regional autonomy. The balance between centralization (harmonization) and peripheralization (disintegration) is fragile, and all agents, regional and central, must practice self-sacrifice and self-restraint to maintain it. A constitution represents a set of rules designed to coordinate the union effectively; the constitution doles out authorities and responsibilities to the various levels of government which both obligates and restricts their actions. Federalism is a network of externalities; the actions taken by one government affect the payoffs to every other.

3.1 The Model

Each government, regional and central, is represented by a single agent, i , $i \in 0, \dots, N$, where 0 represents the central government, and 1 to n the N regions. Agents have finite terms, and face reelection with no term limit. A constitution defines governmental jurisdictions. Actions, such as passing legislation, taxation, or implementing policy, can be considered as a degree of compliance with the constitutional rules, so each agent chooses an action a_i , $a \in [0, 1]$, with $a_i = 1$ representing full compliance with the constitution. Ideally, all agents comply fully with the constitutional rules; the symptoms

¹³For a discussion regarding their effect, see Bednar et al (1995), Bednar & Eskridge (1995), Choper 1977, Eskridge & Ferejohn (1995), Kramer 1995, Weingast 1995.

of instability discussed above all represent non-compliance. The task at hand is to understand the motivation for non-compliance.

Alternatively, agents may pursue policies that deviate from the rules but capture attention and praise from their constituents, such as providing pork, or encroaching upon another level's authority in order to claim credit for favorable programs.¹⁴ Strategic non-compliance can be moderate or extreme, and is represented by values less than 1, and in particular, $a_i \approx 0$ for more extreme non-compliance. Actions are not directly observable, but can be inferred (sometimes imperfectly) by the performance of the union.

Voters have preferences over outcomes alone and reward their agents with reelection based upon the agent's ability to deliver happy outcomes. Therefore, agents are impatient to accomplish popular actions in their terms, represented by a non-zero discount rate, δ . No moral hazard exists between voter and agent: $U_{voter} = U_{agent}$, and U_i may be considered to represent the value of constituency service, or the worth of the agent to the voters. Agents are not rewarded for the accomplishments of others.

Agent i 's utility is represented by the following general form:

$$U_i = U(a_i, a_{-i})$$

where $a_{-i} = (a_1, \dots, a_{i-1}, a_{i+1}, \dots, a_n)$.

To capture the problem of non-compliance, I make the following assumptions.

A1

$$\frac{\partial U_i}{\partial a_i} < 0, \quad \forall i$$

A2

$$\frac{\partial U_i}{\partial a_j} > 0, \quad \forall i, \quad \forall j \neq i$$

A3

$$U_i(\mathbf{0}) < U_i(\mathbf{1}) \quad \text{where} \quad \mathbf{0} = (0_0, \dots, 0_n) \quad \text{and} \quad \mathbf{1} = (1_0, \dots, 1_n)$$

¹⁴See Bednar 1997 for the micromotives of competition between levels of government. In short, politicians compete with one another to claim credit for popular policies, and to distance themselves from unpopular ones, such as taxes.

The first assumption expresses each agent's reluctance to comply with the rules of the federal bargain. According to the second assumption, agents benefit when other agents comply with the union's constitutional provisions. The third assumption eliminates all cases where it is more costly for a member to be a part of the union than out of it; in this model, each agent prefers the case of complete compliance, including its own, to abandoning the union altogether.

Finally, the problem is modelled as an infinitely-repeated game, where a_{it} is the action by agent i in period t . I suppress t when it causes no confusion.

The game proceeds as follows. At the start of each period, each agent chooses a policy a_{it} , which is distorted by idiosyncratic noise, ϵ_i , $E(\epsilon) = 0$, and then generates a value $(a_i + \epsilon_i)$ to the federation. Each agent is able to disburse to its constituents a share of the value of the union β , plus whatever additional reward it gained for itself and its constituents from non-compliance: $(1 - a_i)$. The objective of each agent is to maximize the single period utility:

$$U_i = \beta_i \sum (a_i + \epsilon_i) + (1 - a_i)$$

3.2 Full Information

To derive intuition about strategies, we consider the full information case, where we suppress the random variable ϵ , so the value of the union plainly reveals the action taken by each agent.

I first consider the simplest possible conception of a federal "game" with two agents and a single-dimensional issue space. Each agent chooses an action a_i , $a \in [0, 1]$. Formally, the single period utility function is as follows:

$$\text{Let } U_i = \beta(a_1 + a_2) + (1 - a_i)$$

Assumptions A1 to A3 above continue to hold. Each agent's payoff is the sum of (1) its expected benefit from attainment of the federal objective, with the share represented by a commonly-known parameter (β), set constitutionally, ranging between $1/(n\delta - \delta + 1)$ and 1,¹⁵ where n is the number

¹⁵While this assumption might seem cumbersome, the intuition is easily grasped. Each agent's share of the union's prosperity, β , will always be strictly greater than $\frac{1}{n}$. Proof:

of agents (here, $\beta \in [1/(1 + \delta), 1]$) and (2) the agent's attempt to cull more from the union for the benefit of her constituents.

In this game, unlike the prisoner's dilemma, the action space is continuous, taking any value between 0 and 1, inclusive. In the prisoner's dilemma, the choice space is discrete: players may cooperate or defect. Also, in its general form, the model includes more than two players. However, a comparison to the prisoner's dilemma provides intuition.

In a single-shot prisoner's dilemma, no player wants to be caught complying when the other deviates, but each would like to deviate if the other complies. Although all players prefer the cooperative equilibrium of (comply, comply) to (defect, defect), in the single shot prisoner's dilemma, the unique Nash equilibrium is for both players to deviate, yielding a sub-optimal payoff.

In the repeated prisoner's dilemma, as long as the discount rate is sufficiently low, compliance can be sustained with a tit-for-tat punishment mechanism. Strategies are contingent upon the play in the previous round: each agent cooperates until one defects, at which point they pull the "grim trigger," and defect forever.

With this intuition in mind, we turn to the case of a continuous choice space. While many equilibria exist in the repeated game setting, I will concentrate on the full contribution equilibrium, where $a_i = 1$.

Full contribution can be achieved either with a trigger strategy or with a finite punishment strategy. A trigger strategy, in which deviations are punished forever, is described as:

$$\frac{1}{n\delta - \delta + 1} > \frac{1}{n}$$

Rearranging terms yields:

$$n - 1 > \delta(n - 1)$$

which is true, since $\delta \in (0, 1)$. $\square$

This assumption means that a federation must be greater than the value of the sum of its parts: each agent must get more out of the union than it contributes. The ability of the union to generate prosperity for its members makes the union worth the cost of membership. In short, β gives us the minimum condition for participation in the union.

$$\begin{aligned}
a_{i,t} &= 1 \\
a_{i,t+1} &= 1 \quad \text{if } a_{1,t'} + a_{2,t'} = 2, t' \leq t \\
&= 0 \quad \text{else.}
\end{aligned}$$

NB: With n agents, the strategy is equivalent:

$$\begin{aligned}
a_{i,t} &= 1 \\
a_{i,t+1} &= 1 \quad \text{if } \sum a_{i,t'} = n, \quad \text{for } t' \leq t \\
&= 0 \quad \text{else.}
\end{aligned}$$

PROOF: Suppose Player 1 is playing the above trigger strategy, and Player 2 complies in each period, playing 1. Player 2's payoff would be $2\beta \sum_{t=0}^{\infty} \delta^t$, or $2\beta/(1-\delta)$. Suppose instead that Player 2 chooses to deviate. It suffices to look at period 1 deviations (because with later period deviations, the earlier cooperative periods cancel out). In the first period, Player 2 gets $\beta + 1$, and 1 in each period thereafter. Therefore, Player 2's payoff is $\delta/(1-\delta) + \beta + 1$.

It suffices to show:

$$\frac{2\beta}{1-\delta} > \frac{\delta}{1-\delta} + \beta + 1$$

Rearranging terms yields:

$$\beta > \frac{1}{1+\delta}$$

which holds by assumption. $\square$

Given the assumptions on δ and β , a finite punishment strategy equilibrium also exists. T , the number of periods, will vary inversely with δ and β , the discount parameter and value of the union to each agent, respectively.

$$\begin{aligned}
a_{i,t} &= 1 \\
a_{i,t+1} &= 1 \quad \text{if } a_{1,t} + a_{2,t} = 2 \\
a_{i,t+j} &= 0 \quad \text{else, for } j = 1, \dots, T \text{ periods}
\end{aligned}$$

PROOF: The infinite game full contribution equilibrium payoff and the payoff from deviations can be written as:

$$\begin{aligned} C^\infty &= 2\beta + \delta \frac{2\beta}{1-\delta} \\ D^\infty &= \beta + 1 + \delta \frac{1}{1-\delta} \end{aligned}$$

respectively.

In the finite punishment game, payoffs from cooperating and deviating are, respectively:

$$\begin{aligned} C^T &= 2\beta + 2\delta\beta \frac{1-\delta^{T+1}}{1-\delta} \\ D^T &= \beta + 1 + \delta \frac{1-\delta^{T+1}}{1-\delta} \end{aligned}$$

Which can be rewritten as:

$$\begin{aligned} C^T &= C^\infty - 2\delta\beta \frac{\delta^{T+1}}{1-\delta} \\ D^T &= D^\infty - \delta \frac{\delta^{T+1}}{1-\delta} \end{aligned}$$

If $C^\infty > D^\infty$, then $\exists T^*$ s.t. $C^{T^*} > D^{T^*}$. $\square$

Perfect information imposes strict structure: each agent knows the effort that all others have contributed to the union, and the value of the union to each agent is also known. Under such circumstances full compliance is possible: the federalism achieves the cooperative equilibrium that is efficient and first best; the dual preferences of coordinated unity and local protection are optimally balanced, at least according to the constitution. With no uncertainty about the actions of others, about the ideal balance, or about the value of the union, the federalism will be perfectly stable (or at least no less stable than a unitary state), suffering no domestic disruptions to its harmony.

3.3 Imperfect Information

In a world where future innovations do not threaten by their very uncertainty, and present action is well understood, the incentives to compete with one another might be resolved by hope for long-lived harmonious and profitable relations. But uncertainty, either about the future, or about the present, and especially about the impact of action, complicates monitoring and punishment. The competitive incentives are not totally dispelled, and with the addition of uncertainty, they reemerge to throw off the federal plan, sometimes violently and irretrievably.

In a federalism, uncertainty assumes three distinct roles. First, uncertainty blurs the translation from policy announcement to administrative outcome, the latter being felt by the rest of the union. Random effects, ranging from technological glitches to bad weather, can cause a well-intended policy to go bad.

Furthermore, uncertainty makes it impossible to separate action from exogenous shock, to distinguish between malicious opportunism and random bad luck, because the effort exerted by each agent is not plainly observable. While regional or central agents pass legislation, or announce policy, seeming to lie within its jurisdiction, and these signals are available for all to read and interpret, they do not always accurately indicate the action taken by the agent. An agent might fail to enforce or otherwise distort the administration of policy to produce an effect of non-compliance.

Also, at times an agent will announce a policy that fails to meet the full amount of effort expected of it by the federal bargain. Certain vicissitudes of situation can sometimes excuse such deviations from full compliance. For example, at times it is physically impossible to meet the expectations of the union; more often, political will renders the expected action impossible. Over time, a shift in the public opinion, when universally felt, will cause expectations and the distribution of powers to be altered by convention.

I amend the simple model above to account for uncertainty.

$$\text{Let } U_i = \beta_i(a_1 + a_2 + \epsilon_1 + \epsilon_2) + (1 - a_i),$$

where ϵ_i is a random variable, $E(\epsilon) = 0$. Assume participant i knows ϵ_i after contributing, but not ϵ_{-i} . A similar model has been analyzed by Green and Porter (1984) and Porter (1983) in the context of Cournot competition between firms. The utility maximizing equilibrium is as follows:

$$\begin{aligned}
a_{i,1} &= 1 - \theta \\
a_{i,t+1} &= 1 - \theta \quad \text{if } a_{1,t} + a_{2,t} + \epsilon_{1,t} + \epsilon_{2,t} > \tau \\
a_{i,t+j} &= 0 \quad \text{else, for } j = 1, \dots, K \text{ periods}
\end{aligned}$$

Note the similarity to the perfect information full contribution equilibrium: all agents play a finite-period trigger strategy to ensure compliance. However, the introduction of the uncertainty term causes two changes to the equilibrium strategies. First, since the value of the union to each agent is affected by the uncertainty terms, the trigger amount reflects this uncertainty. The threshold amount, τ , is set exogenously. If the participants see a total contribution that falls below the threshold, they enter a punishment stage. The expected contribution, formerly $a_i = 1$, is also adjusted due to the imperfection of information, and the participants agree to allow a small amount of slippage, represented by θ . Therefore, "full" contribution in this model is represented by $a_i = 1 - \theta$.

The solution concept is represented by the following Markov dynamic programming problem. Let a^* be a candidate for an equilibrium. We can define the present discounted value for agent i if it plays a strategy of a_i in each period and if the others play a^* as follows:

$$\begin{aligned}
V_i(a_i, a_{-i}^*) &= U_i(a_i, a_{-i}^*) + \text{Prob}\left\{\sum_{j \neq i} a_j^* + a_i + \epsilon \geq \tau\right\} \cdot \delta V_i(a_i, a_{-i}^*) + \\
&\quad \text{Prob}\left\{\sum_{j \neq i} a_j^* + a_i + \epsilon < \tau\right\} \cdot \left(\sum_{t=1}^T \delta^t \cdot 1 + \delta^{T+1} \cdot V_i(a_i, a_{-i}^*)\right)
\end{aligned}$$

The vector of strategies a^* is an equilibrium if and only if

$$V_i(a_i^*, a_{-i}^*) \geq V_i(a_i, a_{-i}^*) \quad \text{for all } a_i \in [0, 1]$$

Equivalently,

$$\frac{\partial V_i(a_i, a_{-i}^*)}{\partial a_i} = 0$$

and

$$\frac{\partial^2 V_i(a_i, a_{-i}^*)}{\partial a_i^2} \leq 0$$

In a moment, I will show how certain restrictions on the value function guarantee that $a_i^* < 1$, in which case we can write $a_i^* = 1 - \theta_i$, where θ_i denotes the amount of slippage from full compliance.

The proof of the second equilibrium in this paper follows Green & Porter (1984). Green & Porter demonstrate that under uncertainty, an oligopoly with noncooperative incentives can achieve cooperative behavior by use of periodic punishment regimes. Firms agree to limit their production to drive up prices from the competitive price to one approaching the monopolistic price. Recall that price is determined by both supply and demand; the firms face an uncertain demand, which affects the market price. Firms agree upon a trigger price; when market price falls below the trigger price, they punish by reverting to Cournot competitive production for a limited amount of time. While each has an incentive to deviate from the agreed upon production level, overproduction lowers price; therefore, deviations raise the probability of punishment. In equilibrium, the number of punishment periods and the production level are set to make each firm exactly indifferent between the one period gain by deviation and the present value loss incurred by entering into a punishment regime.

Conceptually, my representation of a federalism mirrors the cartel of Green & Porter, and therefore the proof can be adopted for the present model. I offer a graphic proof below; the technical reader is encouraged to consult Green & Porter (1984) for a formal proof.

An agent's decision to comply or deviate is based upon a comparison of the single-shot benefit to deviation and the present value of the cost of deviation, which is the probability that the deviation is detected multiplied by the loss incurred in the punishment regime. First, consider two situations where costs and benefits of deviations are both linear (that is, when the costs and benefits of deviation rise proportionately to the amount of the deviation). Three possibilities exist: first, costs could consistently outweigh benefits, as they do in Figure A-1. In this case, no amount of deviation is ever worth the cost, and no agents deviate. In the second case, Figure A-2, benefits always outweigh the costs. All agents will always deviate, to the fullest extent. Cooperation is not possible. A third case, not pictured, is when costs and benefits exactly match one another, that is, the two functions are identical. In this case, the result is unpredictable; agents might comply to any degree between zero and one.¹⁶

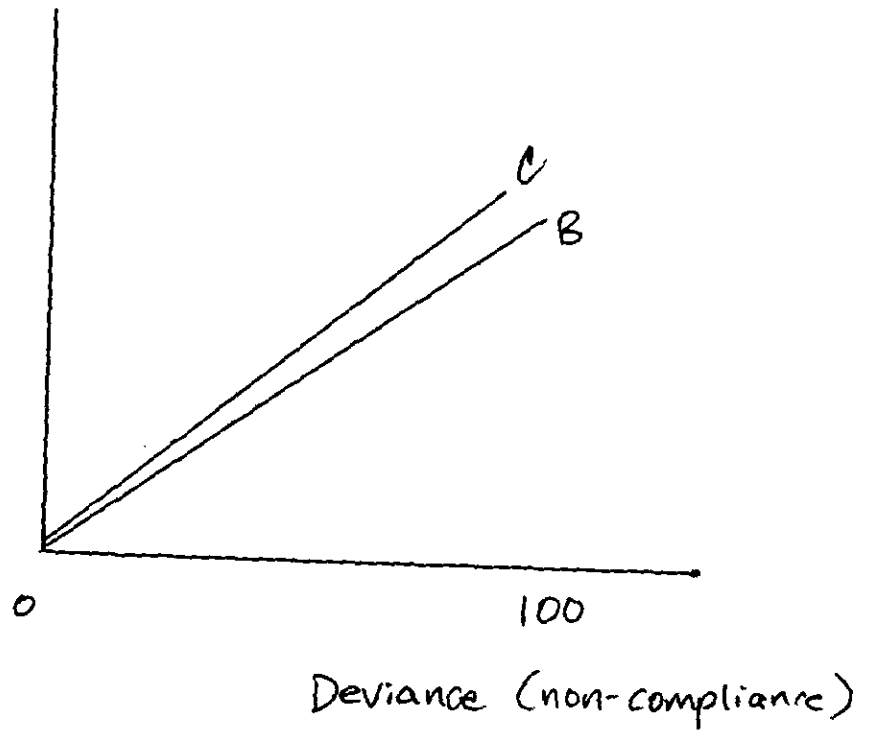
¹⁶To get around the problems generated by the second and third cases, game theorists typically assume that the discount rate is low (agents value the future highly) to drive up

Benefits,
Costs

Figure A-1

Costs outweigh
benefits: No deviance.

$$[\text{Cost} = \text{prob}(\text{caught}) \cdot \text{loss}]$$



Benefits,
Costs

Figure A-2

Benefits outweigh
costs: Total deviance.

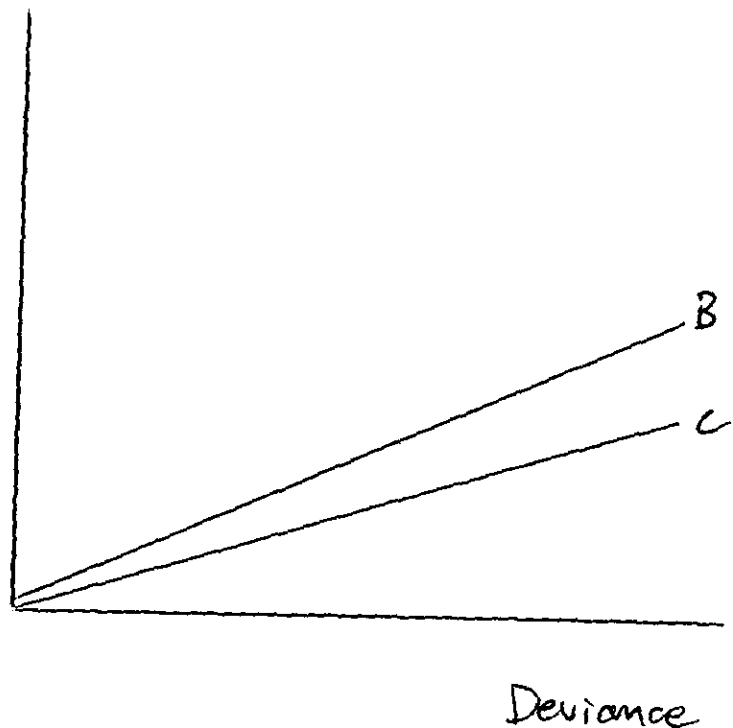


Figure A-3

Concave Benefits:
Moderate Deviance

C, B

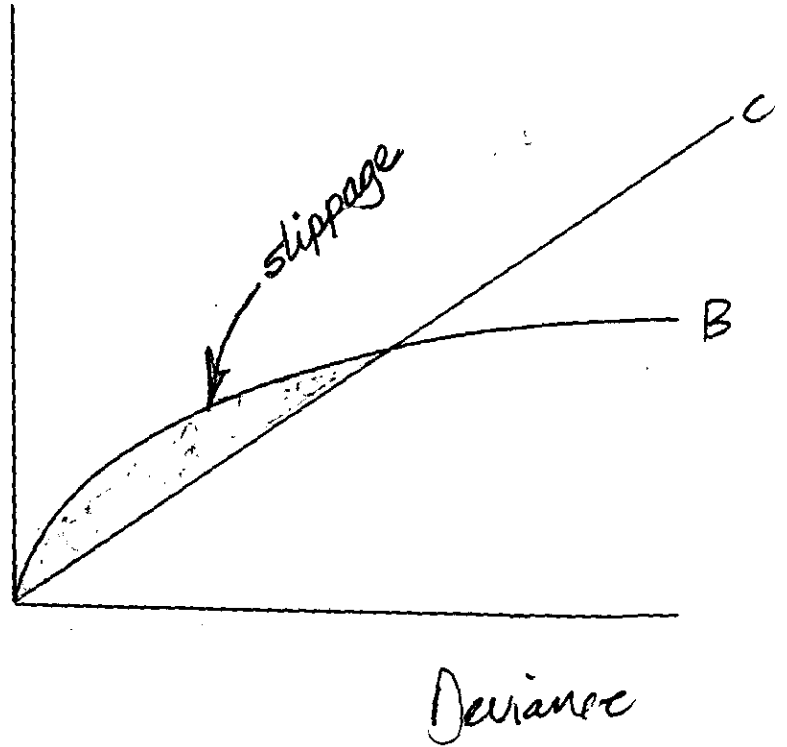


Figure A-4

Convex Costs:
Moderate Deviance

C, B

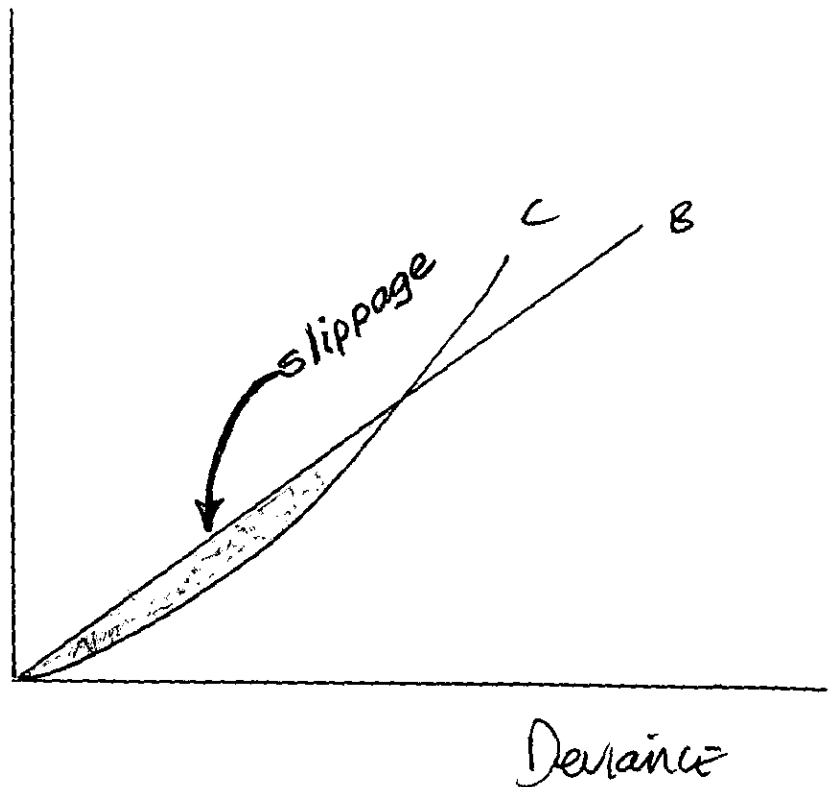


Figure A-5

Concave Benefits
and Convex Costs:
Moderate Deviance

C, B

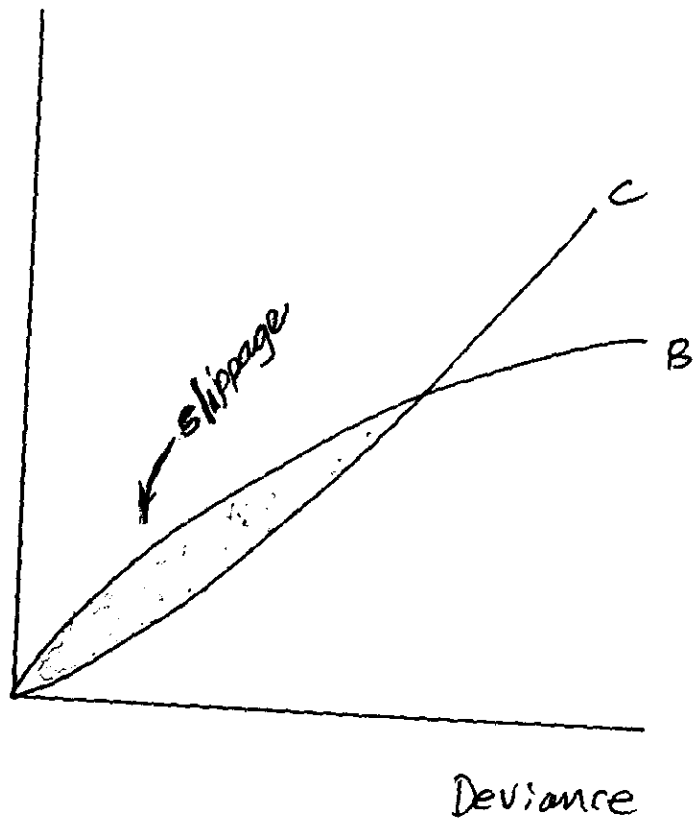
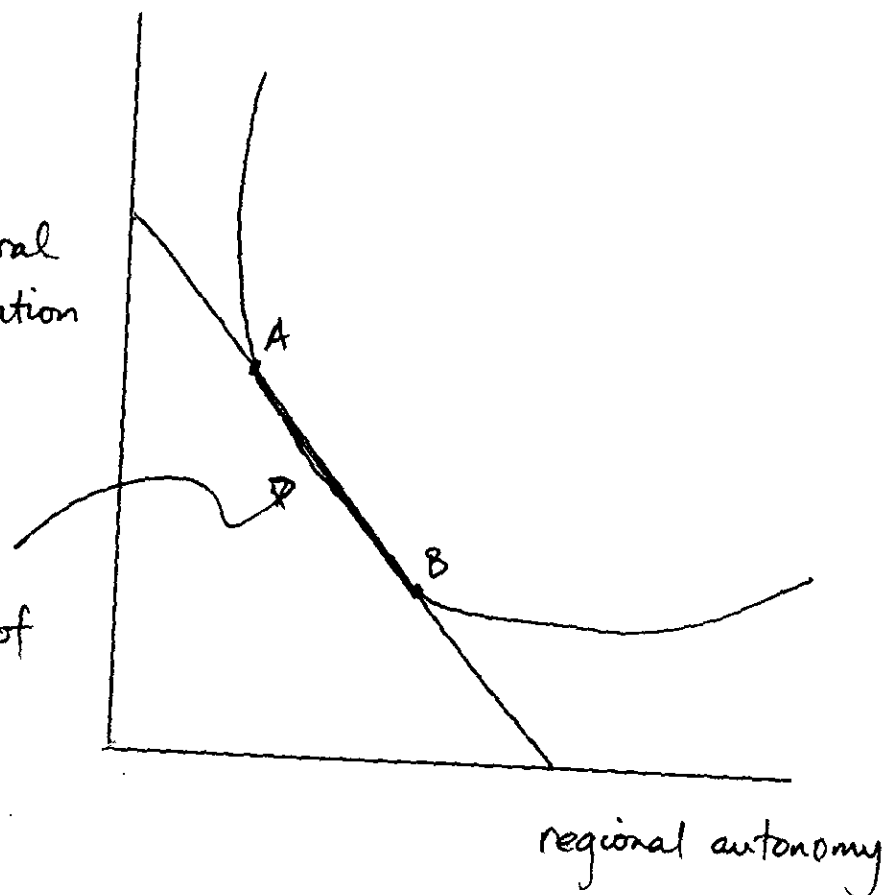


Figure A-6

A flattened citizen
indifference curve,
between points A & B. central
coordination

Constraint:
distribution of
government
authorities



References

- Bednar. 1997. The Credit Assignment Problem. USC Law MS.
- Bednar, Jenna and William N. Eskridge, Jr. 1995. ``Steadying the Court's 'Unsteady Path': A Theory of Judicial Enforcement of Federalism.'' {\it Southern California Law Review} 68:1447-1491.
- Bednar, Jenna, William N. Eskridge, and John Ferejohn. A Political Theory of Federalism. Stanford manuscript.
- Bednar, Jenna, John Ferejohn and Geoff Garrett. A Political Theory of European Federalism. {\it International Review of Law \& Economics} VOL.
- Beer, Samuel H. 1993. {\it To Make a Nation: The Rediscovery of American Federalism.} Cambridge, MA: Belknap Press.
- Brown, Jennifer Gerarda. 1995. ``Competitive Federalism and the Legislative Incentives to Recognize Same-Sex Marriage.'' {\it Southern California Law Review} 68:745-839.
- Caplan, Bryan. 1996. ``Federalism as a Facilitating Practice for Government Cartelization.'' Princeton University Department of Economics manuscript.
- Cremer, Jacques and Thomas R. Palfrey. 1997. ``Political Confederation.'' Caltech manuscript.
- Cremer, J. and T. Palfrey. 1996. ``In or Out?: Centralization by Majority Vote,'' {\it European Economic Review} 40:43-60.
- Dixit, Avisnash and John Londregan. 1996. ``Fiscal Federalism and Redistributive Politics.'' Unpublished manuscript dated July 24, 1996.
- Duchacek, Ivo D. 1970. Comparative Federalism: The Territorial Dimensions of Politics. New York: Holt, Rinehart and Winston.
- Elazar, Daniel. 1987. Exploring Federalism. Tuscaloosa, AL: University of Alabama Press.
- Enrich, Peter D. 1996. ``Saving the States from Themselves: Commerce Clause Constraints on State Tax Incentives for Business,'' {\it Harvard Law Review} 110: 377-468.
- Eskridge and Ferejohn 1995. Vanderbilt Law Review.
- Franck, Thomas M., ed. 1968. Why Federations Fail: An Inquiry into the Requisites for Successful Federalism. New York: New York University Press.

- Friedrich, Carl J. 1968. Trends of Federalism in Theory and Practice. New York: Frederick A. Praeger, Publishers.
- Green, Edward J. and Robert H. Porter. 1984. Noncooperative Collusion under Imperfect Information. *Econometrica* 52(1): 87-100.
- Hicks, Ursula K. 1978. Federalism: Failure and Success, A Comparative Study. New York: Oxford University Press.
- Kramer, Larry. 1994. "Understanding Federalism." *Vanderbilt Law Review*, vol. 47, pp. 1485-1561.
- Lemco, Jonathan. 1991. Political Stability in Federal Governments. New York: Praeger Publishers.
- Madison, ``The Vices of the Political System of the United States,`` in Rutland et al, eds., *The Papers of James Madison*, vol. 9, 1975
- Maggi, Giovanni. Unpub.
- Oates, Wallace E. 1972. *Fiscal Federalism*. New York: Harcourt, Brace, Jovanovich.
- Persson & Tabellini. 1996a. *JPE*.
- Persson & Tabellini. 1996b. *Econometrica*.
- Persson, Torsten, Gerard Roland, and Guido Tabellini. 1996. ``Separation of Powers and Accountability: Towards a Formal Approach to Comparative Politics.`` IGER Working Paper No. 100.
- Peterson, Paul E. 1995. *The Price of Federalism*.
- Porter 1983.
- Riker, William H. 1964. *Federalism: Origin, Operation, Significance*. Boston: Little, Brown and Company.
- The Federalist*, Everyman edition.
- Tiebout, Charles M. 1956. ``A Pure Theory of Local Expenditures.`` *Journal of Political Economy* 64: 416-.
- Tullock, Gordon. 1969. ``Federalism: Problems of Scale.`` *{\it Public Choice}* 19.
- Wechsler, Herbert. 1954. *The Political Safeguards of Federalism: The Role of the States in the Composition and Selection of*

the National Government. {\it Columbia Law Review} 54:543-560.

Weingast, Barry R. 1993. ``The Political Foundations of the Antebellum American Economy.'' Hoover manuscript.

Weingast, Barry R. 1995. ``The Economic Role of Political Institutions: Market-Preserving Federalism and Economic Growth.'' {\it Journal of Law, Economics, and Organization} VOL. pages.

Weingast, Barry R. 1993. ``The Political Foundations of Democracy and the Rule of Law.'' Working paper, IRIS, University of Maryland.

Weingast, Barry R. 1994a. ``Constructing Trust: The Political and Economic Roots of Ethnic and Regional Conflict.'' Stanford manuscript.

Weingast, Barry R. 1994b. {\it Institutions and Political Commitment: A New Political Economy of the American Civil War Era.} Unpublished manuscript, Hoover Institution.