

Renegotiation in Repeated Games

Debraj Ray, October 2006

1. INTRODUCTION

The huge multiplicity of equilibria given by the folk theorem motivates an obvious question: why would players “deliberately” select on equilibria with bad outcomes if “better” equilibria are available? A simple answer to this is that *individual rationality* (along with the common knowledge of the game and strategic beliefs) does not take us further than equilibrium behavior. In particular, it does not permit us to choose among equilibria on the basis of *collective* rationality.

At the same time, we do study repeated games in the hope that it will be able to explain “collusive” outcomes, in which the players can get high payoffs. Thus one justification of the bad equilibria is that we are not really interested in them *per se*, but only insofar as they support the good outcomes. Thus we, the players, might have a conversation about how to proceed and then turn off all conversation, so that when a deviation has to be punished we go ahead and punish, on the expectation that others will too.

But what if we cannot turn off the conversation? Then after a deviation the past is sunk. Might we not have a collective incentive to ignore the deviation and simply start cooperating all over again? We might, but in that case cooperation may not be sustainable to start with! This leads us into the realm of equilibria that are “immune to collective rethinking”, or *renegotiation-proof*.

The following instance illustrates the main point:

Example. Consider the 3×3 two-player game shown below:

	A_2	B_2	C_2
A_1	4, 4	0, 5	0, 0
B_1	5, 0	3, 3	0, 0
C_1	0, 0	0, 0	1, 1

Think of repeating this game once. We know that if β is high enough there is a SGP equilibrium in which we play (A_1, A_2) today, following it up with (B_1, B_2) in the case of compliance and (C_1, C_2) in the case of noncompliance. This is a subgame perfect equilibrium and its equilibrium payoff is the very best one among all equilibria.

Is that good enough grounds to settle on this equilibrium? Maybe, if players can talk just once at the start of the game and not thereafter? But what if they can talk at the start of the repetition as well? Is it then reasonable, using the same collective rationality postulate that started us off on cooperation in the first place, to assume that (C_1, C_2) will ever be played? Probably not.

But now notice the paradox: if it is not reasonable, then we destroy the even more cooperative outcome (A_1, A_2) in the first period! It cannot be sustained any longer.

Renegotiation-proof equilibria are a way to formalize the idea that players can not just talk and cooperate, they can do so repeatedly as the game wears on. As we have just seen, such equilibria do *not* generally give you the best SGP outcome.

First think about a definition for finitely repeated games.

2. A DEFINITION FOR FINITELY REPEATED GAMES

We shall use the familiar device of supporting payoffs to define the concept; it can easily be translated into a definition using strategies (see Bernheim and Ray (1989) for such a definition). We will also restrict ourselves to two-player games from this point on: while the definition of renegotiation proof equilibrium applies, without alterations, to n -player situations, issues of coalitional deviations also become important in such contexts.

Consider a game that is repeated T times, i.e., played $T+1$ times. Use the notation $\omega(A)$ to denote the weak Pareto frontier of some set A : i.e., $\omega(A) \equiv \{p \in A \mid \text{there is no } p' \in A \text{ such that } p' \gg p\}$.

Define S^0 to be the set of one-shot equilibrium payoffs, and let $R^0 \equiv \omega(S^0)$. This is the set of renegotiation-proof payoffs in the one-shot game. There isn't much action here: simply pick the weak Pareto-frontier.¹

Inductively, suppose that we have defined the set R^t as the set of renegotiation-proof payoffs in the t -repeated game. To figure out the R^{t+1} , first consider the set of all payoffs S^{t+1} that can be supported with R^t ,

$$S^{t+1} \equiv \phi(R^t),$$

and then define the set of renegotiation-proof payoffs for the $(t+1)$ -repeated game as

$$R^{t+1} \equiv \omega(S^{t+1}).$$

This completes the recursive definition.

The structure of the sets $\{R^t\}$ can be quite interesting, as the following example demonstrates.

3. AN EXAMPLE

Consider a lender and borrower. The lender lends to the borrower at some fixed rate of interest $r > 0$. There are two projects A and B that the borrower can invest in, yielding net rates of return to the borrower of $\alpha > \beta > 0$. These projects are a matter of complete indifference to the lender, but in any case the lender can dictate the choice of project. At any date, there is a fixed exogenous penalty π for a default of any size on an ongoing loan. Finally, assume that there is an upper bound on the "bad" project B , given by a loan size of $\bar{L}$. On project A , assume no such bound (or a sufficiently larger bound, as the computations below will make clear).

Begin with the stage game. It is clear that there are only two equilibria. To find them, define

$$\ell_0 \equiv \frac{\pi}{1+r}.$$

Now observe that all loan sizes below ℓ_0 will be repaid by the borrower in the stage game. Consequently,

$$S^0 = \{(r\ell_0, \alpha\ell_0), (r\ell_0, \beta\ell_0)\},$$

¹We use the weak Pareto frontier as our criterion in keeping with the idea that every player must strictly wish to renegotiate. This point is brought out clearly in the example below.

and

$$R^0 = \omega(S^0) = V^0.$$

Now turn to the set of all payoffs that can be supported in the game repeated once. (Forget about discounting for this example.) It should then be clear that apart from ℓ_0 , an extra amount of loan can be sustained without fear of default, simply by the lender (credibly) threatening to revert to the bad project in the last period in the case of default in the first period. The (present value) loss to the borrower in that case is given by $(\alpha - \beta)\ell_0$, so that the maximum loan size in the first period of the two-period game is given by

$$\ell_1 \equiv \ell_0 + \frac{(\alpha - \beta)\ell_0}{1 + r}.$$

EXERCISE. If $\ell_1 \leq \bar{L}$, carefully find the value of S^1 , and then show that

$$R^1 = \{(r\ell_1, \alpha\ell_1), (r\ell_1, \beta\ell_1)\}.$$

Recursively, as long as $\ell_t \leq \bar{L}$, we may define in exactly the same way,

$$\ell_{t+1} \equiv \ell_t + \frac{(\alpha - \beta)\ell_t}{1 + r},$$

and then deduce that

$$R^{t+1} = \{(r\ell_{t+1}, \alpha\ell_{t+1}), (r\ell_{t+1}, \beta\ell_{t+1})\},$$

provided that $\ell_{t+1} \leq \bar{L}$. This recursion continues until we reach *first* date T (as we certainly must) such that

$$\ell_T > \bar{L}.$$

At this date, check that R^T must be the singleton set given by

$$R^T = \{(r\ell_T, \alpha\ell_T)\}.$$

If the finite-horizon game has a horizon longer than this, the entire process must build up again from this point! The idea is that in the game repeated T periods, there is exactly *one* renegotiation proof payoff. Consequently, in the game repeated $T + 1$ times, all that can be sustained at the initial date is the original loan size ℓ_0 ! For longer games, the cyclical path builds itself up again, just as outlined above.

It is therefore possible for renegotiation-proof equilibria to exhibit “periodic breakdowns” of cooperation *on* the equilibrium path, and indeed, to select such paths as the unique outcome. This illustrates well the consequences of applying the same selection criteria (in this case, Pareto-optimality) to all subgames as well as on the initial equilibrium path.

4. INFINITELY REPEATED GAMES

Now we turn to a definition of the concept for infinitely repeated games.² Say that a strategy profile σ is *weakly renegotiation proof* (WRP) if it is a SGPE and for all pairs of histories (h_t, h'_s) (where $s = t$ is allowed), the payoff vectors $F(\mathbf{a}(\sigma, h_t))$ and $F(\mathbf{a}(\sigma, h'_s))$ are mutually Pareto-incomparable.

²The same definition appears in Bernheim and Ray (1989) and Farrell and Maskin (1989).

This is easily seen to imply the following feature. Let $P(\sigma)$ be the set of all payoff vectors generated by σ , following all histories. Then $P(\sigma)$ is self-generating, and $\omega(P(\sigma)) = P(\sigma)$.

[Prove this.]

The following observations are relevant.

[1] We could alternatively have taken the feature in the paragraph above to be the defining feature of a WRP set. A WRP set need not be associated with a single equilibrium: it is more like a set of payoffs that has a self-referential consistency property.

[2] This self-referentiality leads to conceptual problems. Observe that the singleton set consisting of the payoff vector generated by any Nash equilibrium of the stage game is WRP. WRP sets are by no means unique.

EXERCISE. Consider the Prisoner's Dilemma given by

2, 2	0, 3
3, 0	1, 1*

Observe that (as discussed) $\{(1, 1)\}$ is a WRP set. By appropriately choosing the discount factor, find another equilibrium that is WRP, and nowhere makes use of the mutual defection cell. Describe precisely the WRP set that it generates.

[3] Thus there is a tension in the "choice" of WRP sets: what is the appropriate theory of the game that players should adopt? One obvious answer is to choose the "best" WRP set: one that is not Pareto-dominated by any other point on any other WRP set. This is a requirement of *external consistency*, as you can tell. The WRP set itself is not just the only criterion that is being used, but a comparison *across* WRP sets is being made.

Unfortunately, the external consistency requirement is not always met. There may not exist *any* WRP set with the required property described in the preceding paragraph. The issue of external consistency then becomes problematic. This is as far as we need to go in this course: see Bernheim and Ray (1989) for a detailed discussion of this and related points.

[4] However, even on the grounds of internal consistency alone, WRP sets are suspect. To see this, consider the following example:

8, 8	0, 0	0, 0
0, 0	0, 0	1, 2*
0, 0	2, 1*	0, 0

Consider the set of payoffs $W \equiv \{(1, 2), (2, 1)\}$. Check that this is indeed WRP.

Now suppose that players indeed hold to W as a theory of how the game will be played from "tomorrow" onwards. In that case, observe that W supports more than W itself: in stages, we see that it covers all the combinations in which $(1, 2)$ and $(2, 1)$ can be played, at the very least. This covers the line segment joining $(1, 2)$ to $(2, 1)$; at least, all the rational convex combinations (weighted by the discount factor) of the two. Thus points approximately halfway between the two extremes become available. But now observe that with such a set, it is possible to sustain the collusive payoff $(8, 8)$ in the first period. This gets us into trouble, because such payoffs

Pareto-dominate segments of W , which will be eliminated as we finally apply the map ω . Thus a truly *internally consistent* renegotiation-proof set may lie pretty far away from W . For more on these issues, see Ray (1994).

With these qualifications in mind, let us return to the study of WRP equilibria. First, a definition. Say that a payoff vector v can be *sustained as a WRP payoff* if there is some WRP equilibrium σ such that $v \in P(\sigma)$.

Following Farrell and Maskin (1989), it is possible under some assumptions to obtain a characterization of WRP payoffs. Let F^{**} be the convex hull of the set of all feasible, *strictly* individually rational payoffs; i.e.,

$$F^{**} \equiv \{v \in F | v \gg 0\}.$$

In what follows, we shall assume that for every $v \in F^{**}$, there is an action vector $a \in A$ such that $f(a) = v$. Recall that for the folk theorem, we also made an assumption like this at the beginning, and then argued after the theorem that such an assumption can be dropped costlessly. This assumption cannot be dropped with equal ease. I will return to this point below.

THEOREM 1. Assume (G.2) and the assumption in the previous paragraph. Let $v \in F^{**}$. Suppose that there are action vectors $a^i \in A$, for $i = 1, 2$ such that

$$(1) \quad \begin{aligned} d_i(a^i) &< v_i \text{ for } i = 1, 2, \\ f_j(a^i) &\geq v_j \text{ for } j \neq i. \end{aligned}$$

Then v is a WRP payoff for all β sufficiently close to unity.

Moreover, if $v \in F^{**}$ is a WRP payoff for any β , then there exist $a^i \in A$, for $i = 1, 2$ such that

$$(2) \quad \begin{aligned} d_i(a^i) &\leq v_i \text{ for } i = 1, 2, \\ f_j(a^i) &\geq v_j \text{ for } j \neq i. \end{aligned}$$

Proof. Sufficiency. Let a be an action vector that attains the payoff v . This is going to be the initial path, while the action vectors a^1 and a^2 are going to serve as punishments.

Begin by observing that there exists $\beta^1 \in (0, 1)$ such that if $\beta \in (\beta^1, 1)$,

$$(3) \quad (1 - \beta)M + \beta d_i(a^i) < v_i,$$

for $i = 1, 2$, where M , it will be recalled, is the value of the maximum absolute payoff in the stage game. Because (3) is strict, there exists a vector p such that for each i ,

$$(4) \quad p_i > d_i(a^i)$$

and

$$(5) \quad (1 - \beta)M + \beta p_i < v_i$$

for all $\beta \in (\beta^1, 1)$.

The idea, now, will be to replicate the value of p_i by playing a^i T times (where T is an integer to be determined), and then go back to the normal phase of playing a . Thus what we want is

$$(6) \quad p_i = (1 - \beta^T)f_i(a^i) + \beta^T v_i$$

for some integer T .

To go about this, let us substitute the RHS of (6) into the inequalities (3) and (4), and see what we need. Note that once we settle on a T , (4) is not going to be a problem for β close enough to unity, because $v_i > d_i(a^i)$ by assumption. For inequality (5) to hold, it must be the case that

$$(1 - \beta)M + \beta[(1 - \beta^T)f_i(a^i) + \beta^T v_i] < v_i$$

must hold for all β close enough to unity. Note that the LHS of the expression above equals v_i at $\beta = 1$. So for the desired result, we need the derivative of the LHS with respect to β to be positive, evaluated at $\beta = 1$. Taking the derivative, we obtain the expression

$$M + f_i(a^i)[1 - (T + 1)\beta^T] + (T + 1)\beta^T v_i,$$

and evaluating this at $\beta = 1$, we get

$$-M - f_i(a^i)T + (T + 1)v_i$$

so that the required condition is

$$(7) \quad (T + 1)v_i > f_i(a^i)T + M.$$

This can be guaranteed for large T , because $v_i > d_i(a^i) \geq f_i(a^i)$. Choose T satisfying (7). Then there is some $\beta^* \in (\beta^1, 1)$ such that if $\beta \in (\beta^*, 1)$, conditions (4), (5) and (6) all hold.

Now define three paths as follows:

$$\begin{aligned} \mathbf{a}^0 &\equiv (a, a, a, \dots), \text{ and} \\ \mathbf{a}^i &= (a^i, a^i, \dots, a^i) \text{ } T \text{ times} \\ &= (a, a, a, \dots) \text{ thereafter,} \end{aligned}$$

for $i = 1, 2$.

We claim that $\sigma(\mathbf{a}^0, \mathbf{a}^1, \mathbf{a}^2)$ is a WRP equilibrium. To establish this, first let's check for subgame perfection.

Deviations from $\mathbf{a}^0$. If i deviates from $\mathbf{a}^0$, he gets *at most*

$$(1 - \beta)M + \beta p_i,$$

by construction. By (5), this is less than v_i .

Deviations from $\mathbf{a}^i$. Suppose, first, that i deviates. The most tempting deviation is in the first period, by the construction of the punishment path. In this case, the total payoff is

$$(1 - \beta)d_i(a^i) + \beta p_i.$$

Using (4), this is less than p_i , so that deviations by i from $\mathbf{a}^i$ are not profitable.

Likewise, j will not deviate from $\mathbf{a}^i$, because by the assumption that $f_j(a^i) \geq v_j$ and the nature of the path $\mathbf{a}^i$, (5) applies right away to prevent deviations (check this).

To complete the proof of sufficiency, all we have to do is check that no two payoff vectors generated by $\sigma(\mathbf{a}^0, \mathbf{a}^1, \mathbf{a}^2)$ ever Pareto-dominate each other. This follows directly from the properties of a^1 and a^2 relative to the payoff vector v .

Remark on the assumption. As in the folk theorem, we can relax the assumption that the payoff vector v is achieved by a pure action. The problem that we now have is to make sure that the various actions which we shall use to intertemporally simulate v do not end up Pareto-dominating each other, or indeed, Pareto-dominating the payoffs. This is not trivial. For an indication of how this task is accomplished in the case where mixed strategies can be observed, see Farrell and Maskin [1989].

Necessity. Suppose that $v \in F^{**}$ is a WRP payoff; i.e., $v \in P(\sigma)$ for some WRP equilibrium σ . We will prove that an action pair a^1 satisfying the required conditions (2) exists. [The proof for a^2 is completely analogous.]

EXERCISE. Assume (G.2). Prove that if v can be supported as a WRP equilibrium, then there exists a WRP σ such that $P(\sigma)$ is compact. [The idea is, as usual, to use a sequential compactness argument. In case you have problems, look at Farrell and Maskin [1989], Lemma 2, 356–357.]

By the exercise, we may assume without loss of generality that σ has a worst continuation equilibrium for player 1. *Choose, among these, the best for player 2.* Let a^1 be the first period action vector on this worst equilibrium path, and let σ^1 denote the continuation equilibrium starting the period after a^1 . Finally, let v^* be the payoff vector when the worst punishment begins.

We will show that a^1 satisfies all the needed conditions.

Clearly, $v_1^* \leq v_1$. We claim that $v_2^* \geq v_2$. Suppose not. Then $v_2^* < v_2$. If $v_1^* < v_1$ as well, we have a contradiction to WRP. So this must mean that $v_1^* = v_1$. But in this case we contradict the choice of the worst equilibrium (it has to maximize player 2's payoff in the class of all equilibria that are worst for player 1). So $v_2^* \geq v_2$, as claimed.

Our next claim is that $f_2(a^1) \geq v_2^*$. Suppose not. Then $f_2(a^1) < v_2^*$. But then $F_2(\sigma^1) > v_2^*$, because σ^1 is the continuation equilibrium. By the WRP requirement, it follows that $F_1(\sigma^1) \leq v_1^*$. But this contradicts, again, our choice of v^* .

We complete the proof by showing that $d_1(a^1) \leq v_1^* \leq v_1$. As noted, $v_1^* \leq v_1$ by construction. To see that $d_1(a^1) \leq v_1^*$, observe that if this were not the case, player 1 could deviate from his punishment by getting $d_1(a^1)$ in the first period, followed by no less than v_1^* . Therefore $d_1(a^1) \leq v_1^* \leq v_1$, and we are done. $\square$