

NON-BAYESIAN LEARNING UNDER IMPRECISE PERCEPTIONS*

Fernando Vega-Redondo**

WP-AD 92-08

* I am grateful to participants in seminars at Harvard University, University of Minnesota, Virginia Polytechnic Institute, and Universitat Autònoma de Barcelona. Luis Corchón, Ignacio Ortuno-Ortín and an anonymous referee also helped me with their very useful comments. Finally, I acknowledge financial support by the CICYT (Spanish Ministry of Education), projects n° PB 89-0294 and n° PS 90-0156.

** Instituto Valenciano de Investigaciones Económicas and University of Alicante.

**Editor: Instituto Valenciano de
Investigaciones Económicas, S.A.**
Primera Edición Diciembre 1992.
ISBN: 84-482-0101-9
Depósito Legal: V-313-1993
Impreso por KEY, S.A., Valencia.
Cardenal Benlloch, 69, 46021-Valencia.
Printed in Spain.

NON-BAYESIAN LEARNING UNDER IMPRECISE PERCEPTIONS

Fernando Vega-Redondo

ABSTRACT

In a preceding companion paper, a static model of individual decision making was proposed that, due to imprecise perceptions, induces simple and inertial behavior at equilibrium ("status-quo optimal") points. This paper addresses two complementary issues. First, it studies the learning dynamics induced by the model and its relationship with the former (static) equilibrium concepts. Second, it characterizes the behavioral implications of the model, comparing them with those derived from standard decision-theoretic frameworks.

1. INTRODUCTION

In a preceding companion paper (Vega-Redondo [1993]), I have studied a static model of individual decision making where the agent is subject to imprecise perceptions. Specifically, she is assumed to have only imprecise perceptions about what outcomes to expect from actions that she has not previously experienced. My primary objective in that paper was to show that, if the idea of imprecise perceptions is suitably formulated, the model is able to rationalize simple and inertial behavior as the result of optimizing behavior.

In this subsequent paper, the purpose is to extend the analysis in the following two complementary directions.

Firstly, I analyze matters in a dynamic scenario where the agent is able to progressively refine her original perceptions through the accumulation of "experience". The learning induced (which need not be strictly Bayesian) yields a wide range of possible limit behavior. On the one hand, if the agent is fully myopic, it coincides with the set of *status-quo optimal* points. (Status-quo optimality is the equilibrium concept under imprecise perceptions used in the static analysis of Vega-Redondo [1993]). On the other hand, if the agent is far-sighted and the underlying framework is sufficiently "regular", the limit set consists of all the *substantively optimal* points (i.e., optimal points under fully precise perceptions). In every other case, limit behavior ranges between these two poles.

Following the dynamic analysis, the second part of the paper turns to the objective of understanding more precisely (in the Revealed-Preference sense) the behavioral implications of the model. Specifically, it carries out an axiomatic comparison of the behavioral implications corresponding to the two polar types of behavior consistent with it (status-quo and substantive optimality). The main conclusion in this respect is that any behavior which can be rationalized as status-quo optimal can be equivalently rationalized as substantively optimal in terms of an acyclic but possibly intransitive preference relation. This result will be interpreted as establishing an

intuitive link between the idea of imprecise perceptions and "incoherent" preferences.

As explained in Vega-Redondo [1993], the present approach has a number of precedents in the literature. Very concisely, the idea that the status quo should play an important role in understanding decision problems is emphasized, for example, in Bewley [1986, 87] or Samuelson & Zeckhauser [1988]. In a somewhat different vein, Kahneman & Tversky [1979] also point to the importance of the status quo as a reference ("framing") point in actual decision problems. In the latter two papers, the reader will find an extensive discussion of a wide variety of different case studies (field or experimental).

In addition to the preceding references, the particular developments of the present paper display further connections to other two important strands of theoretical literature.

In relation with its the learning dynamics, one has the well-established literature on learning and experimentation in a Bayesian context. Noted representatives of it are, for example, the papers of Grossman, Kihlstrom & Mirman [1977] or McLennan [1986]. More recently, the important paper by Aghion, Bolton, Harris & Jullien [1991] has carried out an exhaustive analysis of Bayesian learning under a variety of different scenarios. I have borrowed from this paper some of its theoretical apparatus. It has also inspired some of my results and permitted an important simplification of the original proofs used in previous versions of the paper.

Finally, in relation with the Revealed-Preference analysis of the paper there is the well-established tradition whose initial developments go back to the seminal work of Samuelson [1938]. Following his lead, Houthakker [1950], Richter [1966, 1971] or Moulin [1985] are some of the important further contributions. Some of their results will be part of our future discussion.

The paper is organized as follows. Section 2 includes a concise description of the original static framework which represents the starting point for the present analysis. Section 3 discusses the learning dynamics.

Section 4 focuses on the behavioral characterizations and comparisons. Finally, Section 5 closes the paper with a summary. For the sake of smooth exposition, all proofs are grouped in an Appendix.

2. A BRIEF REVIEW OF THE STATIC MODEL

Since a more detailed description of the model can be found in the companion paper, Vega-Redondo [1993], the account included here is kept specially compact.

Consider a single agent set out to choose some particular alternative in a given set $\mathbb{X} \subseteq \mathbb{R}^m$, her action space. For each $x \in \mathbb{X}$, a corresponding outcome $\omega \in \Omega$ occurs. For simplicity, I shall make $\Omega = [0,1]$ (any completely ordered set would do). Let preferences over outcomes be represented by some von Neuman-Morgensten (VNM) utility function

$$V: \Omega \longrightarrow \mathbb{R}, \tag{1}$$

assumed increasing. If the agent knew accurately the *outcome function*

$$\psi: \mathbb{X} \longrightarrow \Omega, \tag{2}$$

which links actions and outcomes, the decision problem could be formalized as usual, i.e., the agent would be modelled to choose an action that induces a preference-maximal outcome.

The outcome function, however, will not be assumed accurately perceived by the agent. Only the outcome associated to a certain action $x_0 \in \mathbb{X}$, labelled the status quo, is assumed perfectly known (since it has been "experienced"). For any other action $x \neq x_0$, its associated outcome $\psi(x)$ is perceived only imprecisely. Here, *imprecision* is identified with Bayesian subjective uncertainty over the associated outcome. Thus, instead of the particular outcome $\psi(x)$, the agent is assumed to perceive some non-degenerate probability measure around it, as explained below.

The key postulate of the model is that "imprecision grows with 'distance' (or diversity)". In order to model this idea, consider a certain function

$$d: \mathbb{X} \times \mathbb{X} \longrightarrow \mathbb{R}_+ \quad (3)$$

which measures diversity between alternatives. (It is assumed that, for any $x, x' \in \mathbb{X}$, $d(x, x') = d(x', x)$ and $d(x, x') = 0$ iff $x = x'$.)

Associated to any given $x \in \mathbb{X}$, consider a family of perceptions associated to it for each of the particular distances $d \in \mathbb{R}_+$ from which it can be evaluated (i.e., distances to the status quo). Since $\Omega = [0, 1]$ is totally ordered, these perceptions may be represented by a corresponding family of cumulative distribution functions

$$\{F(x, d)\}_{d \in \mathbb{R}_+}. \quad (4)$$

In line with the previous discussion, we assume that, for $d = 0$ and all $x \in \mathbb{X}$,

$$F(x, 0)(\omega) = 0, \text{ if } \omega < \psi(x); \quad (5a)$$

$$F(x, 0)(\omega) = 1, \text{ if } \omega \geq \psi(x). \quad (5b)$$

That is, at zero distance (i.e., when x itself is the status quo) the perception is fully concentrated in $\psi(x)$.

The additional idea that *increasing distance brings about increasing imprecision* is formalized through the notion of Second Order Stochastic (SOS) dominance, applied to the family of perceptions (4) associated to any given action x . Specifically, it is assumed that, for all $x \in \mathbb{X}$, if $d_1 > d_2$, then $F(x, d_2)$ strictly SOS-dominates $F(x, d_1)$. That is:

$$\int_{\omega=0}^{\omega'} F(x, d_1)(\omega) d\omega \geq \int_{\omega=0}^{\omega'} F(x, d_2)(\omega) d\omega, \quad (6)$$

for all $\omega' \in \Omega$, with the inequality holding strictly for some $\hat{\omega} \in \Omega$.

It is important to understand the motivation underlying such modelling assumption. Intuitively, one good like to be sure that, however one formalizes imprecision, the more imprecise a perception is the worse off any "reasonable" agent is. Thus, suppose the agent is risk averse (i.e., V is strictly concave). Then, the concept of SOS-dominance precisely formalizes this idea.

For, as it is well-known (see Fishburn & Vickson (1978)), given two prospects (or perceptions), one is SOS-dominated by the other if, and only if, the latter is preferred to the former by any risk-averse agent.

The basic features of the model can be summarized through the following assumptions:

(RA) Risk Aversion: The VNM utility function V in (1) is continuous, increasing, and strictly concave.

(SA) Status-Quo Accuracy: For all $x \in \mathbb{X}$, $F(x, 0)$ is given by (5).

(II) Increasing Imprecision: For all $x \in \mathbb{X}$, $F(x, \cdot)$ satisfies (6).

Within this scenario, the agent is supposed to confront the following decision problem: Choose a perceived maximal action within a certain choice set $\Gamma \subseteq \mathbb{X}$. The following definition states the equilibrium condition characterizing status-quo (or simply, q-)optimality.

Definition 1: An action $x_0 \in \mathbb{X}$ is **q-optimal** if $x = x_0$ maximizes $U_{x_0}(x) \equiv \int V(\cdot) dF(x, d(x_0, x))$ in the set Γ (with respect to x).

One of the main concerns of the paper is to compare the behavior captured by the previous definition with that resulting from the standard decision-theoretic approach with precise perceptions. The latter shall be labelled substantive (or s-) optimality. For completeness, its definition closes this section.

Definition 2: An action $x_0 \in \mathbb{X}$ is **s-optimal** if $x = x_0$ maximizes $V(\psi(x))$ in the set Γ (with respect to x).

3. LEARNING DYNAMICS

Consider an agent who must adopt an action x_t at each one of a series of dates, $t = 0, 1, 2, \dots$. Throughout these dates, her choice set $\Gamma \subseteq \mathbb{X}$ remains fixed. So happens with the von Neuman-Morgensten utility function $V: \Omega \rightarrow \mathbb{R}$

which represents, for all t , her current ("instantaneous") preferences over outcomes $\omega \in \Omega$.

As above, the agent will be assumed, in general, uncertain over the particular function that links actions to outcomes. Denote by Ξ the set of continuous outcome functions of the form $\psi: \mathbb{X} \rightarrow \Omega$ and endow it with the supremum norm. The set of possible agent's perceptions will be identified with the set $\Lambda \equiv \Delta(\Xi)$ of Borel probability measures over Ξ . At each t , some $\lambda_t \in \Lambda$ formalizes her current subjective uncertainty over the underlying outcome function $\hat{\psi} \in \Xi$, known to remain unchanged throughout.

Living in an intertemporal context, the agent views her decision problem as an intertemporal one.

On the one hand, her preferences over any given stream of outcomes $(\omega_t)_{t=1}^{\infty}$ are assumed given by the corresponding discounted payoffs $\sum_{t=1}^{\infty} \delta^{t-1} V(\omega_t)$, with $0 \leq \delta < 1$. As before (c.f. (RA) in Section 2), the function $V(\cdot)$ is assumed continuous, increasing, and strictly concave. Moreover, the choice set Γ is assumed compact.

On the other hand, the agent takes into account the learning possibilities entailed by any planned path of actions. Learning, here, will not necessarily be restricted to being of the Bayesian type (see below). It is generally formalized through a certain learning mechanism, which is defined as follows.

Definition 3: A Learning Mechanism (LM) is a mapping $L: \Lambda \times \mathbb{X} \times \Xi \rightarrow \Lambda$ which, for each probability measure $\lambda \in \Lambda$, action $x \in \mathbb{X}$, and underlying outcome function $\hat{\psi} \in \Xi$, yields a revised probability measure $\lambda' = L(\lambda, x, \hat{\psi})$.

Of course, in order for a given LM to represent a sensible learning process some minimal properties must be required from it. I now postulate some of them.

First, it should be the case that, given any series of realized observations $(x_t, \hat{\psi}(x_t))_{t=0,1,\dots,t_0}$, the agent's perceptions must:

- (i) not contradict any of them;
- (ii) assign some positive probability to the "true" underlying outcome function $\hat{\psi}$.

Formally, I shall require the following:

(A.1) Learning from experience

Let $\{x_t, \lambda_t\}_{t=0,1,\dots}$ be a sequence of actions and perceptions satisfying $\lambda_{t+1} = L(\lambda_t, x_{t+1}, \hat{\psi})$, $t = 0, 1, \dots$. Then, for every t :

- (i) $\text{supp}(\lambda_t) \subseteq \{\psi \in \Xi: \psi(x_\tau) = \hat{\psi}(x_\tau), \tau = 0, 1, \dots, t\}$.
- (ii) $\forall t, \hat{\psi} \in \text{supp}(\lambda_t)$.

The next assumption embodies a two-fold requirement. The first part is an assumption of *tabula rasa*. It specifies that, at the start of the process, initial perceptions reflect only the information induced by the status quo, as formalized by the static model of Vega-Redondo [1993] which is described in Section 2. The second part requires that, along any sequence of further actions, the corresponding perceptions are never less "precise" than those induced by the preceding action alone. Here, the specific meaning of "precision" is formalized, again as described in Section 2, through the concept of SOS-dominance.

To formalize these matters, let $\phi(x, \lambda)$ represent, in analogy with $F(x, d)$, the distribution function over outcomes perceived for action $x \in \mathbb{X}$ when λ prevails. That is:

$$\phi(x, \lambda)(\omega) = \text{Prob}_\lambda \{ \psi(x) \leq \omega \}. \quad (7)$$

I postulate:

(A.2) Learning improves "basic" perceptions

Let $F(\cdot)$, $d(\cdot)$ be as in Section 2 and consider a sequence of actions and perceptions $\{x_t, \lambda_t\}_{t=0,1,\dots}$ satisfying $\lambda_{t+1} = L(\lambda_t, x_{t+1}, \hat{\psi})$, $t \geq 0$. Then

- (i) $\forall x \in \mathbb{X}, \phi(x, \lambda_0) = F(x, d(x_0, x))$.
- (ii) $\forall t, \forall x \in \mathbb{X}, \phi(x, \lambda_t)$ weakly⁽¹⁾ SOS-dominates $F(x, d(x_t, x))$.

Let Λ be endowed with the topology of weak convergence. (Of course, Ξ and $\mathbb{X}$ are endowed with the topology induced by their corresponding metrics, the supremum norm and the Euclidean metric respectively.) The following assumption is also made:

(A.3) Continuity

The function $L: \Lambda \times \mathbb{X} \times \Xi \longrightarrow \Lambda$ is continuous.

In the strictly Bayesian context of Aghion *et al* (1991), the continuity of the learning mechanism requires, in general, the introduction of some noise in the process of observation. In the present context, however, Bayes Rule must be complemented with additional considerations (for example, learning from "close actions") which justify the above assumption of continuity. For the sake of smooth discussion, I postpone to Remark 1 below a clarification of these comments. There, the issue of continuity will be discussed in conjunction with the relationship of the model with standard Bayesian updating.

The next (and last) assumption captures the following intuitive idea: at each stage of the process, the learning induced by past observations should exhaust all possibilities. In other words, everything that can be learned from previous experience has to be already incorporated into the current perceptions. This must imply that the stochastic process $\{\lambda_t\}$ induced by any sequence of actions should be a Martingale, in an appropriate sense, when the perceptions themselves are viewed as a random variable. Formally, this is specified as follows:

(A.4) Exhaustive learning

For all $\lambda \in \Lambda$ and $x \in \mathbb{X}$,

¹ Weak SOS-dominance does not require that the inequality (6) be satisfied strictly for some $\hat{\omega} \in \Omega$. This allows for the possibility that x_t be the only relevant source of information underlying λ_t .

$$E_{\lambda} \left[\int f dL(\lambda, x, \cdot) \right] = \int f d\lambda,$$

for any bounded and continuous real function $f: \Xi \rightarrow \mathbb{R}$.

Assumption (A.4) states that, given current perceptions λ and any planned action x , the expected perception after the observation (as induced by $L(\cdot)$) is "indistinguishable close" (according to the assumed weak topology on Λ) to the current perception. This assumption would be satisfied if L reflected, for example, Bayesian learning. But, in general, it seems an indispensable requirement of any model of learning. If it were violated, the agent would be able to refine her perceptions without any further observation.

In Example 1 below, I shall discuss a particular LM which satisfies assumptions (A.1) to (A.4). Besides illustrating matters, this example will also establish the logical consistency of all our postulated assumptions.

We are now in a position to describe the agent's intertemporal decision problem: Given initial conditions $(x_0, \lambda_0) \in \mathbb{X} \times \Lambda$,

$$\text{Maximize } E_{\lambda_0} \left[\sum_{t=1}^{\infty} \delta^{t-1} V(\omega_t) \right], \quad (P)$$

with respect to decision rules $\{g_t\}_{t=1}^{\infty}: \Lambda \rightarrow \mathbb{X}$, where for all $t = 1, 2, \dots$:

- (i) $x_t = g_t(\lambda_{t-1}) \in \Gamma$;
- (ii) $\omega_t = \hat{\psi}(x_t)$;
- (iii) $\lambda_t = L(\lambda_{t-1}, x_t, \hat{\psi})$.

The above items do not require much explanation. Point (i) expresses that the decision rule has to be measurable with respect to the available information and prescribe a feasible action. Point (ii) reflects the link between actions and outcomes established by the underlying (and imprecisely perceived) outcome function $\hat{\psi}$. Finally, point (iii) indicates that the agent is able to take into account her future learning opportunities.

This section includes two basic results. The first one (Theorem 1 below) establishes the existence of an optimal action path and the convergence of the

learning process. The second result (Theorem 2) specifies the range of possible limit behavior under alternative further assumptions.

Theorem 1: *The decision problem (P) has an optimal (and stationary) decision rule $g^*: \Lambda \rightarrow \mathbb{X}$. Moreover, for every such rule, there is some $\tilde{\lambda}^* \in \Lambda$ such the path of induced perceptions $\{\lambda_t^*\}_{t=1}^\infty \rightarrow \tilde{\lambda}^*$.*

Proof: See the Appendix.

Given the underlying outcome function $\hat{\psi}$, let $\eta^*(x_0)$ denote the set of limit points consistent with initial conditions x_0 and some associated λ_0 which satisfies (A.2). (That is, for all $x \in \mathbb{X}$, $\phi(x, \lambda_0)$ weakly SOS-dominates $F(x, d(x_0, x))$.)⁽²⁾ Correspondingly, let

$$\eta^*(\Gamma) \equiv \bigcup_{x_0 \in \Gamma} \eta^*(x_0). \quad (8)$$

Finally, denote by X^* and X^{**} the set of q-optimal points and s-optimal points, respectively, as described in Definitions 1 and 2 above. As shown in Vega-Redondo [1993], $X^{**} \subseteq X^*$, the inclusion being strict under natural assumptions. In the present dynamic context, the following result holds.

Consider the following regularity condition:

(R) (i) For all $\psi \in \Xi$, the function $V(\psi(\cdot)): \mathbb{X} \rightarrow \mathbb{R}$ is continuously differentiable and quasi-concave. (ii) Γ is a smooth manifold.

Theorem 2: *The set $\eta^*(\Gamma) \subseteq X^*$. Moreover:*

(i) *If $\delta = 0$, $\eta^*(\Gamma) = X^*$.*

(ii) *If $\delta > 0$ and (R) holds, then $\eta^*(\Gamma) = X^{**}$.*

² Even though λ_0 is a sufficient state variable for the dynamical system, (A.2) requires that it satisfy the above SOS-dominance condition for some "initial action" x_0 . Thus, varying x_0 , one traces the full range of allowed initial perceptions.

Proof: See the Appendix.

Theorem 1 is a basic existence result. It establishes that the learning model proposed is non-vacuous, and always yields a well-defined limit amount of learning. On the other hand, Theorem 2 bears upon the scope of different behavior which is consistent with the model. It shows that, depending on initial conditions, such behavior ranges from:

- (i) the whole span of q-optimal behavior if the agent is fully myopic (i.e., completely discounts the future), to
- (ii) the set of fully optimal points (i.e., s-optimal) if the underlying features of the model are sufficiently regular (smooth and quasi-concave).

The intuition behind the inclusion $\eta^*(\Gamma) \subseteq X^*$ in Theorem 2 is easy to explain. In the limit, when all learning possibilities are exhausted (see Theorem 1), intertemporally optimal behavior also becomes myopically optimal. By (A.2), it must also be q-optimal, i.e., optimal under "basic" imprecise perceptions. This shows the inclusion $\eta^*(\Gamma) \subseteq X^*$. The equality $\eta^*(\Gamma) = X^*$ when $\delta = 0$ is then obvious since, under our assumptions, every q-optimal point may remain stationary in an optimal ("myopic") path.

As for part (ii) of Theorem 2, the intuitive argument is based on the combination of the following two facts. First, since $V(\psi(\cdot))$ is assumed smooth, local exploration on the gradient of this function can be conducted at arbitrarily small cost. Therefore, no limit point can be locally sub-optimal (in the substantive sense). Combining this fact with the assumed quasi-concavity of the previous function, the global optimality follows.

Remark 1: Bayes Rule and learning

In our context, learning is not just the outcome of statistical updating (or Bayes Rule) based on the observation of realized outcomes. It should also incorporate the additional information (in general, only imprecise) which any chosen action provides about neighboring actions. (Recall the original motivation of the model in Section 2.) To clarify this point, it will be useful to focus our discussion on the particular issue of continuity in learning, which was raised above in relation with assumption (A.2).

As pointed out by Aghion *et al* [1991], the application of Bayes Rule alone will in some cases lead to discontinuous learning. To illustrate this fact, consider the following simple context. Suppose that the prevailing subjective probability measure of the agent over the underlying outcome function is concentrated on, say, just two (continuous) functions: ψ_1 and ψ_2 . Assume that these functions only coincide (i.e., cross) at some action $\hat{x}$, yielding different outcomes elsewhere. Thus, if the agent observes the outcome associated to any action different from $\hat{x}$, no matter how close to it, the agent obtains full information. However, at $\hat{x}$ itself, her prior is not refined at all. Learning, therefore, is discontinuous at $\hat{x}$.

The previous example notwithstanding, continuity of learning has been assumed by (A.2). This must imply, in particular, that we must be ready to accept Bayes-Rule violations in some cases. To fix ideas, consider again the previous schematic context. In it, continuity of learning must imply (since, by (A.1), the outcome associated to any adopted action is assumed to be observed accurately) that the agent learns fully by doing any action (even action $\hat{x}$). Thus, Bayes Rule is obviously violated. How can this be justified ?

As suggested above, the justification relies on the fact that, in our context, the adoption of an action should not be conceived as simply yielding a "point observation" that reveals its associated outcome. The action chosen should also convey ("uncertain") information on other actions; the closer these are (as reflected by the diversity function $d(\cdot)$), the better the entailed information. (Recall (SA) and (II) in Section 2.)

Thus, in the above example, the choice of action $\hat{x}$ should be viewed as giving arbitrarily precise information about any action which is sufficiently close. Since full information obtains from any of these, it seems natural to postulate that the same will follow from doing action $\hat{x}$. This is the main rationale behind the assumption of continuity of learning postulated by (A.2). As explained, its possible incompatibility with Bayes Rule should not be viewed as implying some degree of "incoherent thinking". Rather, it should be attributed to the different nature (in particular, information content) of what represents an "observation" in our framework.

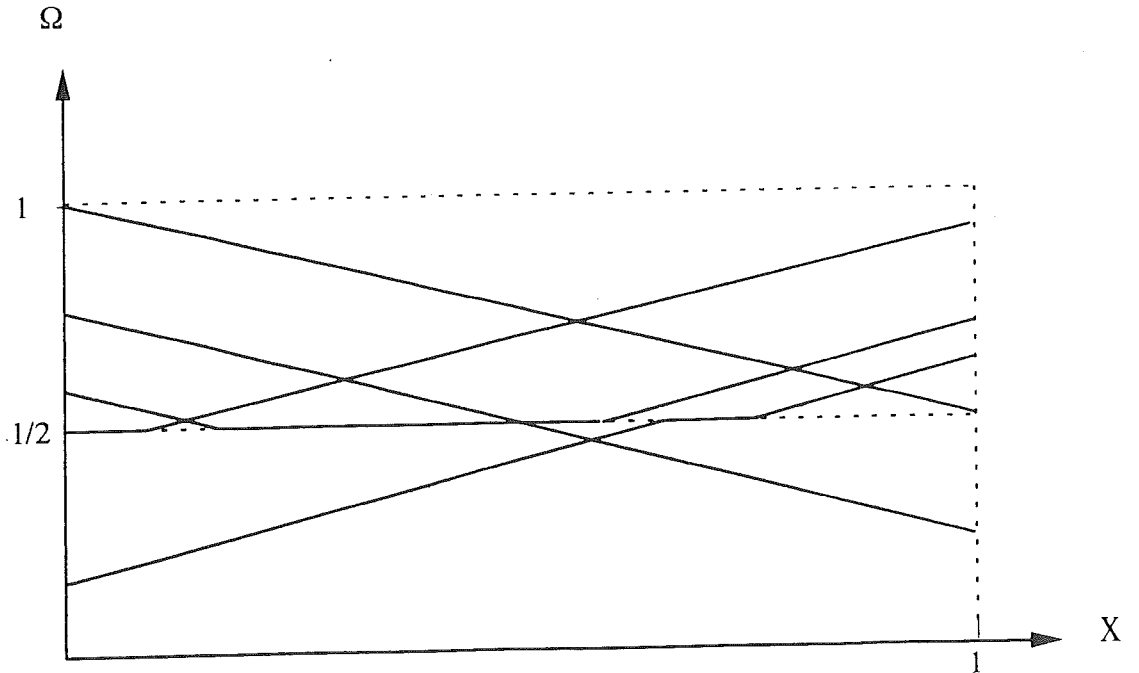
Example 1: A Learning Mechanism satisfying (A.1) to (A.4)

Make $\mathbb{X} = [0,1]$. Consider any pair of actions $(x', x'') \in \mathbb{X}^2$ with $x' \leq x''$, and a pair $(a, b) \in \{1, -1\}^2$. Associated to $z \equiv [(x', x''), (a, b)]$, an outcome function $\psi_z: \mathbb{X} \rightarrow \Omega (= [0,1])$ is defined as follows:

$$\begin{aligned} \forall x \in [x', x''], \psi_z(x) &= 1/2; \\ \forall x < x', \psi_z(x) &= 1/2 + a/2 (x' - x); \\ \forall x > x'', \psi_z(x) &= 1/2 + b/2 (x - x''). \end{aligned} \quad (9)$$

Consider the family of outcome functions of the previous type generated by the set $Z \equiv \{[(x', x''), (a, b)] \in \mathbb{X}^2 \times \{1, -1\}^2: x' \leq x''\}$. It is denoted by $\tilde{\Xi}$. This family will represent the *a priori* set of outcome functions considered by the agent. Some functions in $\tilde{\Xi}$ are represented in the following diagram.

FIGURE 1



Denote by $\tilde{\Lambda}$ the set of Borel probability measures on $\tilde{\Xi}$ and let

$$B: \tilde{\Lambda} \times \mathbb{X} \times \tilde{\Xi} \longrightarrow \tilde{\Lambda} \quad (10)$$

stand for the function which, given a prior $\lambda \in \tilde{\Lambda}$, some action $x \in \mathbb{X}$, and the observed outcome $\psi(x) \in \Omega$, induces (when well defined) the posterior $\lambda' = B(\lambda, x, \psi)$ determined by Bayes Rule. As explained in Remark 1, $B(\cdot)$ will behave discontinuously at some action points. Specifically, denote by $z(\psi) \equiv [(x'(\psi), x''(\psi)), (a(\psi), b(\psi))]$ the vector of parameters which characterizes some given ψ . Then, if $z(\psi) \in \text{int } Z$ and λ assigns positive probability mass to outcome functions ψ' with $x''(\psi') > x''(\psi)$, $B(\lambda, \cdot, z)$ is discontinuous "from the left" at actions $x''(\psi)$.

Motivated by the already explained idea that, in our context, the experimentation of some action allows the agent to get arbitrarily accurate information of sufficiently close actions, the aforementioned discontinuities may be smoothed as follows.

First, define a continuous function:

$$\alpha: \mathbb{X} \times \mathbb{X} \longrightarrow [0,1], \quad (11)$$

which satisfies:⁽³⁾

- (i) $1 \geq \alpha(\tilde{x}, \hat{x}) > 0$ if $\tilde{x} < \hat{x}$;
- (ii) $\alpha(\tilde{x}, \hat{x}) = 0$ if $\tilde{x} \geq \hat{x}$.

Given any $\lambda \in \tilde{\Lambda}$, define:

$$\sigma(\lambda) \equiv \{x \in \mathbb{X}: \forall \psi, \psi' \in \text{supp}(\lambda), \psi(x) = \psi(x')\}, \quad (12)$$

i.e., the set of actions whose associated outcome can be ensured with λ -probability one. Note, for future reference, that the particular form of the outcome functions in $\tilde{\Xi}$ (and in particular their continuity) implies that, for all $\lambda \in \tilde{\Xi}$, $\sigma(\lambda)$ is a closed and connected set. Its maximum and minimum will be denoted, respectively, by $\bar{x}(\lambda)$ and $\underline{x}(\lambda)$.

³ Generally, one would like to think of the value $\alpha(\tilde{x}, \hat{x})$ in (i) as $\hat{\alpha}$ -non-decreasing in a suitable notion of "distance" between $\tilde{x}$ and $\hat{x}$.

Denote by

$$\hat{\Lambda} = \{ \lambda \in \tilde{\Lambda}: \sigma(\lambda) \neq \emptyset \} \quad (13)$$

the set of perceptions which can ensure the outcome of at least one action. Any initial perceptions λ_0 which satisfy (A.2(i)) must belong to $\hat{\Lambda}$. (See below for our choice of λ_0 .) Thus, if the learning mechanism proposed,

$$L: \tilde{\Lambda} \times \mathbb{X} \times \tilde{\Xi} \longrightarrow \tilde{\Lambda}, \quad (14)$$

turns out to be closed in $\hat{\Lambda}$ (i.e., $L(\hat{\Lambda}) \subseteq \hat{\Lambda}$), it will be enough to define it on this set alone.

Denote by $v'(x, \lambda)$ and $v''(x, \lambda)$ the posterior probabilities induced by Bayes rule when the prior λ is refined with the information that the underlying outcome function ψ satisfies $x'(\psi) = x$ or $x''(\psi) = x$, respectively. Given any $\lambda \in \hat{\Lambda}$, $x \in \mathbb{X}$, and $\psi \in \tilde{\Xi}$, $L(\lambda, x, \psi)$ is defined as follows:

- (a) If $x \in \sigma(\lambda)$, $L(\lambda, x, \psi) = B(\lambda, x, \psi) = \lambda$.
- (b) If $x > \bar{x}(\lambda)$,

$$L(\lambda, x, \psi) = \alpha(x, x''(\psi)) B(\lambda, x, \psi) + [1 - \alpha(x, x''(\psi))] v''(x''(\psi), \lambda). \quad (15)$$
- (c) If $x < \underline{x}(\lambda)$,

$$L(\lambda, x, \psi) = \alpha(x'(\psi), x) B(\lambda, x, \psi) + [1 - \alpha(x'(\psi), x)] v'(x'(\psi), \lambda).$$

If well-defined, (15) implies that $L(\cdot)$ is closed in $\hat{\Lambda}$, as desired. On the other hand, this formulation is well-defined only if the posteriors $B(\cdot)$, $v'(\cdot)$, and $v''(\cdot)$ are themselves well-defined applications of Bayes Rule. For this to happen, it is enough that the underlying outcome function have positive λ -density, i.e., $\psi \in \text{supp}(\lambda)$. But this will hold throughout since, again, $L(\cdot)$ is closed with respect to this property and λ_0 will be chosen to satisfy it.

The discussion of the example is completed by showing that, under appropriate initial conditions, the postulated LM satisfies all required assumptions (A.1) to (A.4).

Assume, for simplicity, that the underlying outcome function $\hat{\psi} \in \tilde{\Xi}$ is the constant $\hat{\psi}(\cdot) = 1/2$. Choose any initial status quo $x_0 \in \mathbb{X}$ and, denoting by $\lambda^u (\in \tilde{\Lambda})$ the uniform probability measure on $\tilde{\Xi}$, let $\lambda_0 = B(\lambda^u, x_0, \hat{\psi}) \in \hat{\Lambda}$.⁽⁴⁾ Given (15), such choice of initial conditions (x_0, λ_0) obviously implies (A.1). As for (A.2), it is satisfied if we make $d(\cdot)$ equal to the Euclidean distance ($d(x, y) = |x - y|$) and, for each $x \in \mathbb{X}$, $F(x, d)$ corresponds to the uniform distribution on the interval $[1/2 - d/2, 1/2 + d/2]$.

Assumption (A.3) is satisfied, by construction, for all (λ, x, ψ) such that $\lambda \in \hat{\Lambda}$ and $\psi \in \text{supp}(\lambda)$. For our purposes (i.e., ensuring the existence of an optimal decision rule), this is all what matters. However, $L(\cdot)$ could be extended continuously, if desired, to its whole domain $\tilde{\Lambda} \times \mathbb{X} \times \tilde{\Xi}$.

Finally, it must be verified that the martingale property embodied by assumption (A.4) applies to the learning mechanism described by (15). This is trivially true for case (a). For cases (b) and (c), it is easily seen to follow as well since the learning mechanism that is being postulated represents a convex combination of two rules ($B(\cdot)$ and either $v'(\cdot)$ or $v''(\cdot)$), both of which satisfy the required martingale property. (Notice that each of these rules results from the application of Bayes Rule to two different "observations".)

4. REVEALED PREFERENCE

Theorem 2 above specifies different particular scenarios for which the span of limit behavior induced by the learning model ranges from:

(i) the set of q-optimal points, if the agent is fully myopic (or impatient), to

⁴ If x_0 is not "too close" to any of the extreme points of the interval $[0, 1]$ (note that $x'(\hat{\psi}) = 0$, $x''(\hat{\psi}) = 1$), one may think that λ_0 is obtained from applying to λ^u an extension of $L(\cdot)$ to general perceptions in $\hat{\Lambda}$. In this case, it may be assumed that $\alpha_{(x'(\hat{\psi}), x)} = \alpha_{(x, x''(\hat{\psi}))} = 1$ and $L(\cdot)$ and $B(\cdot)$ coincide.

(ii) the (no larger) set of s-optimal points, when the agent is not utterly impatient and the underlying framework is sufficiently regular.

These two polar types of behavior reflect the alternative concepts of optimality which underlie the static analysis conducted in Vega-Redondo [1993]. It is, therefore, of significant interest to have a comparative understanding of the behavioral implications which characterize each of them. This is the main purpose of the present section, which adopts for this purpose the customary Revealed-Preference approach.

The primitive object of analysis is a *pattern of behavior*, namely, a given collection of behavior traced across a certain family of different decision problems. Different decision problems differ in their corresponding choice sets. Preferences, on the other hand, are assumed constant for all them. In this context, the standard Revealed-Preference question is posed as follows: What global features (restrictions) must a given pattern of behavior satisfy in order to be rationalized by (i.e., be consistent with) a certain pre-determined decision-theoretic model ?

To address this question formally let the given family of decision problems under consideration be parameterized by some $\theta \in \Theta$. Associated to each θ , denote the corresponding choice set by $\Gamma(\theta) \subseteq \mathbb{X}$. A precise definition of the concept of pattern of behavior follows.

Definition 4: A **Pattern of Behavior** is a mapping $\eta: \Theta \longrightarrow 2^{\mathbb{X}}$ such that for all $\theta \in \Theta$, $\eta(\theta) \subseteq \Gamma(\theta)$.

Thus, a pattern of behavior η is simply a collection of "feasible" behavior. If the focus is on conditions which may rationalize it as s-optimal, there is a large amount of literature that, starting with Samuelson [1935] and Houthakker [1950], has addressed such revealed-preference question. Abstracting from representability issues,⁵ the precise question becomes

⁵ The model requires that the underlying preference relation over outcomes be representable by some utility function $V(\cdot)$. Thus, we would need some standard separability conditions in addition to the C-Axiom below to ensure that such a representation is possible. Since these issues pose no special problems but are nevertheless quite alien to our main focus, I choose to abstract from them. The interested reader is referred to Fishburn (1970).

whether there is a preference relation $\mathbf{R}$ on $\mathbb{X}$ such that:

$$\forall \theta \in \Theta, \eta(\theta) = \{x \in \Gamma(\theta): x \mathbf{R} y \text{ for all } y \in \Gamma(\theta)\}. \quad (16)$$

As shown by Richter (1966), a behavioral axiom which is both necessary and sufficient for a pattern of behavior η to satisfy (9) in terms of a reflexive, transitive, and total preference relation $\mathbf{R}$ (I shall call such a preference *regular*) is the so-called *Congruence (or C-) Axiom*. To define it formally, a binary relation $\mathbf{Q}$ on $\mathbb{X}$ is first constructed as follows:

$$\forall x, y \in \mathbb{X}, x \mathbf{Q} y \Leftrightarrow \exists \theta \in \Theta: x \in \eta(\theta) \ \& \ y \in \Gamma(\theta). \quad (17)$$

Let T denote the *transitive closure* of $\mathbf{Q}$. The C-Axiom is specified as follows:

$$\begin{aligned} \text{Congruence (C-)Axiom: } & \forall \theta \in \Theta, \forall x, y \in \mathbb{X}, \\ & [x \in \eta(\theta), y \in \Gamma(\theta), y T x] \Rightarrow y \in \eta(\theta). \end{aligned}$$

A pattern of behavior which satisfies (9) for a regular preference relation will be called *s-rational*. Analogously, a pattern of behavior will be called *q-rational* when it can be viewed as *q-optimal* (c.f. Definition 1) for a suitable "perception structure". More precisely:

Definition 5: A pattern of behavior η is said to be **q-rational** if there exist a VNM utility function $V(\cdot)$ satisfying (RA), and a diversity function $d(\cdot)$, and a family of perceptions $\{F(x, d)\}_{(x, d) \in \mathbb{X} \times \mathbb{R}_+}$ satisfying (SA) and (II) such that, for all $\theta \in \Theta$:

$$\eta(\theta) = \{x \in \Gamma(\theta): \int V(\cdot) dF(x, 0) \geq \int V(\cdot) dF(y, d(x, y)) \text{ for all } y \in \Gamma(\theta)\}.$$

Since the concept of *q-optimality* is weaker than that of *s-optimality* (i.e., every *s-optimal* point is also *q-optimal* for any set of perceptions), it is clear that every *s-rational* pattern of behavior can be rationalized as *q-optimal*. As we shall show below, the converse is, in general, not true.

In addressing the revealed-preference characterization of *q-rationality*, one could aim directly at a behavioral axiom analogous to the C-Axiom which permitted a clear comparison with it. It seems more interesting, however, to

do it indirectly by investigating the precise level of "coherence" which must be required from a (status-quo-free) preference relation R in order to be behaviorally equivalent to q -rationality. This approach has the advantage of allowing for a more ready and clear-cut comparison between the status-quo-dependent approach and the standard status-quo-free context. Moreover, as a by-product, one may then readily use the well-established characterization results developed for the status-quo-free context under varying "degrees of coherence" in order to obtain, if desired, behavioral axioms characterizing q -optimal behavior (c.f. Remark 2 below). Our main result in this section reads as follows:

Theorem 3: *A stationary pattern of behavior η is q -rational if, and only if, it is rationalizable, in the sense of (9) above, by an acyclic and weakly representable preference relation.⁽⁶⁾*

In some sense, the precedent result can be seen as identifying the precise extent of incoherent or "incongruent" behavior (that is, behavior violating the Congruence Axiom) which one can "rationalize" in terms of the model proposed. Within the framework of this model, such incoherent behavior (for example, behavior that exhibits intransitive indifference) is not seen as the outcome of an incoherent preference relation. Rather, as the consequence of (explicitly modelled) imprecise perceptions which are experienced by an otherwise perfectly "coherent" agent.

Remark 2: Behavioral axioms characterizing q -rational behavior

If X is finite and $\Gamma(\Theta) = 2^X$ there exist two behavioral axioms -the Chernoff and Expansion Axioms- which are known to characterize behavior rationalizable by an acyclic preference relation. (See, for instance, Moulin

⁶ A preference relation R is called acyclic if its induced asymmetric part P exhibits no cycles. ($\forall x, y \in X, xPy \iff xRy \ \& \ \neg(xRy)$.) The relation P is weakly representable if there exists a weak homomorphism $\overline{w}$ from (X, P) to $(\mathbb{R}, >)$, i.e., if $\forall x, y \in X, xPy \Rightarrow \overline{w}(x) > \overline{w}(y)$.

Note that there is a redundancy in the statement of the Theorem since a weakly representable relation is obviously acyclic. This redundancy is maintained for expositional purposes. As mentioned in Footnote 5, the representability requirement could be replaced by an appropriate separability assumption.

(1985)). By Theorem 3, they also characterize q-rational behavior in this case. (Note that, if $\mathbb{X}$ is finite, the weak representability of the acyclic preference relation is trivially achieved).

I close this section with the discussion of two simple examples which illustrate the behavioral restrictions imposed by q-rationality. The first one shows that the C-Axiom is too strong to be consistent in general with q-rationality. The second indicates that the very weak Q-Axiom (see below for a precise definition) is indeed too weak to ensure q-rationality. These examples are redundant, in view of our precedent discussion (Theorem 3 and Remark 2). I include them, however, since they might be helpful in illustrating the basic features of the model.

Example 2: q-rational behavior need not satisfy the C-Axiom:

Consider a decision problem defined as follows:

- The action set $\mathbb{X} = \{\alpha, \beta, \gamma\}$,
- The outcome space $\Omega = [0, 1]$,
- The VNM utility function $V(\cdot)$ given by $V(\omega) = \omega^{1/2}$ for each $\omega \in \Omega$,
- The set $\Theta = \{\theta_1, \theta_2\}$ which parameterizes two different problems with choice sets $\Gamma(\theta_1) = \{\alpha, \beta\}$, $\Gamma(\theta_2) = \mathbb{X} = \{\alpha, \beta, \gamma\}$,
- The outcome function ψ , defined by $\psi(\beta) = \psi(\gamma) = 1/2$, $\psi(\alpha) = 2/5$.
- A family of perceptions which, in particular, induce from action α as status quo the following stochastic outcome functions for actions β and γ :

(i) $F(\beta, d(\alpha, \beta))$ induces a mean-preserving spread over $\psi(\beta)$ that assigns an equal probability of $1/3$ to the outcomes $\omega_0 = 1/2$, $\omega_1 = 0$ and $\omega_2 = 1$.

(ii) $F(\gamma, d(\alpha, \gamma))$ is a mean-preserving spread over $\psi(\gamma)$ that assigns an equal probability of $1/3$ to the outcomes $\omega_0 = 1/2$, $\omega_3 = 1/4$, and $\omega_4 = 3/4$.

(According to our interpretation of the model, such perceptions would be rationalized by saying that the diversity between α and β is larger than the one between α and γ).

Under those conditions, it is immediate to check that if $\eta: \Theta \longrightarrow \mathbb{X}$ denotes the corresponding pattern of q-rational behavior, we have:

$$\eta(\theta_1) = \{\alpha, \beta\} \text{ and } \eta(\theta_2) = \{\beta, \gamma\}, \quad (18)$$

which clearly violates the C-Axiom.

Example 3: The Q-Axiom need not imply q-rationality

Given a pattern of behavior η and the associated binary relation Q as defined in (17), consider the following axiom:

$$\text{Q-Axiom: } \forall \theta \in \Theta, \forall x \in \mathbb{X}, [x \in \Gamma(\theta) \ \& \ (\forall y \in \Gamma(\theta), x \ Q \ y)] \Rightarrow x \in \eta(\theta).$$

It can be shown (Richter (1971)) that the Q-Axiom characterizes all patterns of behavior that can be rationalized (in the sense of (9)) by some preference relation. In this case, no particular requirement (reflexivity, transitivity of totality) is demanded from the preference ordering. The Q-Axiom is, therefore, a minimal requirement of rationality. As illustrated by the following example, it is too weak to be consistent in general with the concept of q-rationality.

Let the action set $\mathbb{X} = \{\alpha, \beta, \gamma, \chi\}$. For the set of parameters $\Theta = \{\theta_1, \theta_2\}$, let their associated choice sets be $\Gamma(\theta_1) = \{\alpha, \gamma, \chi\}$ and $\Gamma(\theta_2) = \{\beta, \gamma, \chi\}$. Consider the pattern of behavior $\eta: \Theta \longrightarrow \mathbb{X}$ defined as follows: $\eta(\theta_1) = \{\gamma\}$; $\eta(\theta_2) = \{\chi\}$. It is clear that η satisfies the Q-Axiom: any preference relation that includes the binary relation Q induced by η rationalizes in the sense of (9) such pattern of behavior. (In this example, moreover, such preference relation can be chosen complete, v.g. choose R as follows: $\alpha \ R \ \chi \ R \ \gamma \ R \ \alpha \ R \ \beta \ R \ \gamma \ R \ \chi \ R \ \beta$).

The pattern η , however, is not q-rational. If $\psi: \mathbb{X} \longrightarrow \Omega$ and $V: \Omega \longrightarrow \mathbb{R}$ denote the outcome and VNM utility functions which could rationalize η as q-optimal, since $\eta(\theta_1) = \{\gamma\}$, $V(\psi(\gamma))$ must be strictly greater than both $V(\psi(\alpha))$ and $V(\psi(\chi))$. But then, no matter what perceptions consistent with ψ we consider, γ should belong to $\Gamma(\theta_2)$. Not being this the case, η cannot be q-rational.

APPENDIX

Proof of Theorem 1:

Consider first the existence of an optimal decision rule for (P). It is enough to show the existence of a well-defined and continuous value function

$$W^*: \Lambda \longrightarrow \mathbb{R}, \quad (19)$$

which satisfies the Bellman equation:

$$W^*(\lambda) = \max_{x \in \Gamma} \int [V(\psi(x)) + \delta W^*(L(\lambda, x, \psi))] d\lambda(\psi), \quad (20)$$

for all $\lambda \in \Lambda$. Or, equivalently, to show the existence of a function W^* which is a fixed point of the mapping

$$T: \mathbb{C} \longrightarrow \mathbb{R}^\Lambda, \quad (21)$$

where $\mathbb{C} (\subset \mathbb{R}^\Lambda)$ denotes the set of continuous real functions on Λ and, for each $W \in \mathbb{C}$,

$$T(W) = \max_{x \in \Gamma} \int [V(\psi(x)) + \delta W(L(\lambda, x, \psi))] d\lambda(\psi). \quad (22)$$

Endowed with the supremum norm, the set $\mathbb{C}$ is a complete space. Thus, the existence of the required fixed point is a consequence of the following two considerations. Firstly, $T(\mathbb{C}) \subseteq \mathbb{C}$, by the continuity of $V(\cdot)$ and $L(\cdot)$. Secondly, the mapping $T(\cdot)$ is a contraction. To verify this latter point is a quite standard exercise which, for the sake of completeness, we now sketch.

Let W_1 and W_2 be any two functions in $\mathbb{C}$ and denote:

$$\hat{W}_i(x, \lambda) \equiv \int [V(\psi(x)) + \delta W_i(L(\lambda, x, \psi))] d\lambda(\psi), \quad i = 1, 2. \quad (23)$$

Let

$$x_i(\lambda) \in \arg \max_{x \in \Gamma} \hat{W}_i(x, \lambda), \quad (24)$$

and make $x(\lambda)$ an action $x_{i_0}(\lambda)$, $i_0 \in \{1, 2\}$, which satisfies $\hat{W}_{i_0}(x, \lambda) = \max \{\hat{W}_1(x, \lambda), \hat{W}_2(x, \lambda)\}$. Obviously:

$$\begin{aligned}
|T(W_1(\lambda)) - T(W_2(\lambda))| &\leq |\hat{W}_1(x(\lambda), \lambda) - \hat{W}_2(x(\lambda), \lambda)| \\
&= |\delta \int [W_1(L(\lambda, x(\lambda), \psi)) - W_2(L(\lambda, x(\lambda), \psi))] d\lambda(\psi)| \\
&\leq \delta \int |W_1(L(\lambda, x(\lambda), \psi)) - W_2(L(\lambda, x(\lambda), \psi))| d\lambda(\psi) \\
&\leq \delta \|W_1 - W_2\|_\infty,
\end{aligned} \tag{25}$$

where $\|\cdot\|_\infty$ denotes the supremum norm in $\mathbb{C}$. Thus,

$$\sup_{\lambda \in \Lambda} |T(W_1(\lambda)) - T(W_2(\lambda))| \equiv \|T(W_1) - T(W_2)\|_\infty \leq \delta \|W_1 - W_2\|_\infty \tag{26}$$

which proves that $T(\cdot)$ is a contraction in $\mathbb{C}$, since $\delta < 1$. This completes the proof of the first part of the theorem.

The second part of the theorem (the convergence of optimal learning) relies essentially on assumption (A.4). Given a decision rule g (not necessarily optimal), any bounded and continuous real function $f: \Xi \rightarrow \mathbb{R}$ induces, by (A.4), a martingale $\{y_t\}$ with

$$y_t = \int f d\lambda_t, \tag{27}$$

where (λ_t) is the sequence of perceptions induced by L , g , the underlying $\hat{\psi}$, and the given initial conditions (x_0, λ_0) . That is, the stochastic process $\{y_t\}$ satisfies:

$$E[|y_t|] < \infty; \tag{28}$$

$$E[y_{t+1} | y_0, \dots, y_t] = y_t; \tag{29}$$

for all $t = 1, 2, \dots$. By standard convergence theorems on martingales (see, for example, Karlin & Taylor [1975, Theorem 5.1, p. 278]), we know that, since (28) holds uniformly for all t , $y_t \rightarrow \tilde{y}$ for some real $\tilde{y}$, almost surely (a.s.).

The previous conclusion holds for all bounded and continuous real function f . Thus, there must exist some $\tilde{\lambda} \in \Lambda$ such that, in the topology of weak convergence, $\lambda_t \rightarrow \tilde{\lambda}$, a.s. Since, by (A.2(ii)), $d\lambda_t(\hat{\psi}) > 0$, where $\hat{\psi}$ is the underlying outcome function, the proviso "almost surely" can be dispensed with. This completes the proof of the second part of the theorem. ■

Proof of Theorem 2:

The first step of the argument requires showing that if the stochastic process $\{\lambda_t^*\}$ defines a martingale, then $\{W^*(\lambda_t^*)\}$ is a sub-martingale. This conclusion is based on the following preliminary lemma.

Lemma: The value function W^* defined in (20) is convex.

Proof: Fix any decision rule $g: \Lambda \rightarrow \Gamma$ and denote by $U(g; \lambda)$ the expected discounted payoff of following it when the prevailing perception is λ . The function $U(\cdot; \cdot)$ is obviously linear in the second argument. This, as it is now argued, implies the desired convexity of W^* .

For suppose otherwise. Let g^* be an optimal decision rule (whose existence is guaranteed by Theorem 1). By definition, $W^*(\lambda) = U(g^*; \lambda)$ for all $\lambda \in \Lambda$. If W^* is not convex, there exist some $\lambda_1, \lambda_2 \in \Lambda$, $0 < \alpha < 1$, such that:

$$W^*(\alpha\lambda_1 + (1-\alpha)\lambda_2) < \alpha W^*(\lambda_1) + (1-\alpha) W^*(\lambda_2), \quad (30)$$

or

$$U(g^*; \alpha\lambda_1 + (1-\alpha)\lambda_2) < \alpha U(g^*; \lambda_1) + (1-\alpha) U(g^*; \lambda_2), \quad (31)$$

a contradiction with the linearity of $U(g^*, \cdot)$.

Given the convexity of W^* and the martingale properties (28) and (29), Jensen's inequality implies that the real stochastic process $\{W^*(\lambda_t^*)\}$ is a bounded sub-martingale with respect to the stochastic process $\{\lambda_t^*\}$. That is:

$$E[W^*(\lambda_{t+1}^*) | \lambda_0^*, \dots, \lambda_t^*] \leq W^*(\lambda_t^*). \quad (32)$$

Which, by a standard argument in the theory of martingales, implies:

$$E[W^*(\lambda_{t+1}^*) | \lambda_0^*, \dots, \lambda_t^*] - W^*(\lambda_t^*) \rightarrow 0, \text{ a.s.}, \quad (33)$$

when $t \rightarrow \infty$. As above, the proviso "almost surely" can be dispensed with because of (A.2(ii)).

After some algebraic manipulations, rewrite (20) as follows:

$$W^*(\lambda_t^*) - \frac{1}{1-\delta} \int V(\psi(g^*(\lambda_t^*))) d\lambda_t^*(\psi) = \frac{1}{1-\delta} [E[W^*(\lambda_{t+1}^*)] - W^*(\lambda_t^*)]. \quad (34)$$

Taking limits in the above expression, and using (33) and the continuity of W^* , it follows that:

$$\frac{1}{1-\delta} \int V(\psi(g^*(\lambda_t^*))) d\lambda_t^*(\psi) \rightarrow W^*(\tilde{\lambda}^*), \quad (35)$$

where $\tilde{\lambda}^*$ is the limit perception established by Theorem 1. That is, the "capitalized" current payoff converges to the maximum discounted payoff given by W^* .

Define now, for each $\lambda \in \Lambda$,

$$\varpi(\lambda) = \max_{x \in \Gamma} \int V(\psi(x)) d\lambda(\psi), \quad (36)$$

as the maximum expected current payoff which the agent can obtain given current perceptions λ . By the definition of $\varpi(\cdot)$, we have:

$$\varpi(\lambda) \geq \int V(\psi(g^*(\lambda))) d\lambda(\psi), \quad (37)$$

and, by the definition of $W^*(\cdot)$:

$$\frac{1}{1-\delta} \varpi(\lambda) \leq W^*(\lambda), \quad (38)$$

Combining (37) and (38), it follows that:

$$0 \leq W^*(\lambda_t^*) - \frac{1}{1-\delta} \varpi(\lambda_t^*) \leq W^*(\lambda_t^*) - \frac{1}{1-\delta} \int V(\psi(g^*(\lambda_t^*))) d\lambda_t^*(\psi), \quad (39)$$

for all t . And, therefore, by (35):

$$W^*(\lambda_t^*) - \frac{1}{1-\delta} \varpi(\lambda_t^*) \rightarrow 0, \quad (40)$$

as $t \rightarrow \infty$.⁽⁷⁾ Or, by the continuity of W^* and $\bar{w}$,

$$W^*(\tilde{\lambda}^*) = \frac{1}{1-\delta} \bar{w}(\tilde{\lambda}^*), \quad (41)$$

where, as above, $\tilde{\lambda}^*$ denotes the limit perception established by Theorem 1.

Combining (35) and (41), it follows that every limit point $\tilde{x}$ of the sequence $\{g^*(\lambda_t^*)\}$ must maximize current expected payoff for the limit perception $\tilde{\lambda}^*$, i.e., it must satisfy:

$$\forall x \in \Gamma, \int V(\psi(\tilde{x})) d\tilde{\lambda}^*(\psi) \geq \int V(\psi(x)) d\tilde{\lambda}^*(\psi), \quad (42)$$

From assumption (A.2), $\phi(x, \lambda_t^*)$ SOS-dominates $F(x, d(g(\lambda_{t-1}^*), x))$ for all t . Since $\phi(x, \cdot)$ is continuous in λ , the limit point $\tilde{x}$ which satisfies (42) also satisfies:

$$\forall x \in \Gamma, \int V(\psi(\tilde{x})) d\tilde{\lambda}^*(\psi) = V(\hat{\psi}(\tilde{x})) \geq \int V(\psi(x)) dF(x, d(\tilde{x}, x)), \quad (43)$$

i.e., it is q -optimal (c.f. Definition 1). This shows the inclusion $\eta^*(\Gamma) \subseteq X^*$.

The equality $\eta^*(\Gamma) = X^*$ when $\delta = 0$ is an immediate consequence of the fact that every q -optimal point is an "instantaneous" maximal at $t = 0$ since, by (A.2(i)), $\phi(x, \lambda_0) = F(x, d(x_0, x))$ for all $x \in \mathbb{X}$. Thus, if $\delta = 0$, it is also a stationary point and, *a fortiori*, a limit point of the process.

The last part of the theorem, namely the equality $\eta^*(\Gamma) = X^{**}$ when $\delta > 0$ and (R) holds, is proven as follows. Suppose, for the sake of contradiction, that there is a limit point $\tilde{x}$ of $\{g^*(\lambda_t^*)\}$ such that $\tilde{x} \notin X^{**}$, i.e., is not s -optimal. Assume, for simplicity, that this point is interior to Γ .⁽⁸⁾ By (R) it must be that:

$$\frac{\partial V(\psi(\tilde{x}))}{\partial x_{i_0}} \neq 0, \text{ for some } i_0 \in \{1, 2, \dots, m\}. \quad (44)$$

⁷ An analogous conclusion can be found in Theorem A.4 of Aghion et al [1991].

⁸ Otherwise, one simply needs to modify what follows by taking into account the standard boundary conditions on any non-vanishing derivative.

Suppose the process were to start with the initial perception $\tilde{\lambda}^*$, and thereafter the agent follows a decision rule which prescribes action $\tilde{x}$ indefinitely. From the previous considerations, such decision rule must be optimal under these circumstances, yielding an expected discounted payoff equal to:

$$W^*(\tilde{\lambda}^*) = \frac{V(\psi(\tilde{x}))}{1-\delta}. \quad (45)$$

Following Aghion *et al* [1991, Theorem 3.2],⁽⁹⁾ consider the following alternative decision rule:

At $t=1$, choose $\tilde{x}$;

At $t=2$, choose $\tilde{x} + \delta_{i_0}(\epsilon)$ where $\delta_{i_0}(\epsilon) \equiv (0, \dots, 0, \epsilon, 0, \dots, 0)$ and the value of ϵ , which occupies the i_0 -th component of the vector, is chosen small enough such that $\tilde{x} + \delta_{i_0}(\epsilon) \in \Gamma$;

At $t=3$,

(i) if $V(\psi(\tilde{x} + \delta_{i_0}(\epsilon))) \leq V(\psi(\tilde{x}))$, choose $\tilde{x}$ at $t = 3$ and forever after;

(ii) if $V(\psi(\tilde{x} + \delta_{i_0}(\epsilon))) > V(\psi(\tilde{x}))$, choose $\tilde{x} + \delta_{i_0}(\alpha\epsilon)$ at $t = 3$ and forever after, where α is chosen small enough such that, given ϵ , the latter action belongs to Γ .

Denote by $[w]^+ \equiv \max\{w, 0\}$ for any $w \in \mathbb{R}$. The expected payoff of the above described decision rule may be written as follows:

$$\begin{aligned} \zeta(\epsilon, \alpha; \tilde{\lambda}^*) &\equiv V(\psi(\tilde{x})) + \delta \int V(\psi(\tilde{x} + \delta_{i_0}(\epsilon))) d\tilde{\lambda}^*(\psi) \\ &+ \delta^2 \int V(\psi(\tilde{x} + \delta_{i_0}(\alpha\epsilon))) [V(\psi(\tilde{x} + \delta_{i_0}(\epsilon))) - V(\psi(\tilde{x}))]^+ d\tilde{\lambda}^*(\psi) \quad (46) \\ &+ \delta^2 W^*(\tilde{\lambda}^*) \int [V(\psi(\tilde{x})) - V(\psi(\tilde{x} + \delta_{i_0}(\epsilon)))]^+ d\tilde{\lambda}^*(\psi). \end{aligned}$$

Thus, combining (45) and (46), the expected payoff difference of playing the alternative strategy is:

⁹ The line of argument used here is a very slight and straightforward adaptation to a multidimensional action space of the one used by these authors. I reproduce it here, for the sake of completeness.

$$\begin{aligned}
W^*(\tilde{\lambda}^*) - \zeta(\varepsilon, \alpha; \tilde{\lambda}^*) &= \delta \int [V(\psi(\tilde{x} + \delta_{i_0}(\varepsilon))) - V(\psi(\tilde{x}))] d\tilde{\lambda}^*(\psi) \\
&+ \delta^2 \int [V(\psi(\tilde{x} + \delta_{i_0}(\alpha\varepsilon))) - V(\psi(\tilde{x}))] [V(\psi(\tilde{x} + \delta_{i_0}(\varepsilon))) - V(\psi(\tilde{x}))]^+ d\tilde{\lambda}^*(\psi),
\end{aligned} \tag{47}$$

which, dividing by ε and making it tend to zero, becomes:

$$\begin{aligned}
\lim_{\varepsilon \rightarrow 0} \frac{W^*(\tilde{\lambda}^*) - \zeta(\varepsilon, \alpha; \tilde{\lambda}^*)}{\varepsilon} &= \\
&= \delta \int \frac{\partial V(\psi(\tilde{x}))}{\partial x_{i_0}} d\tilde{\lambda}^*(\psi) + \alpha \delta^2 \int \left[\frac{\partial V(\psi(\tilde{x}))}{\partial x_{i_0}} \right]^+ d\tilde{\lambda}^*(\psi).
\end{aligned} \tag{48}$$

By (A.1) and (A.5), the second term of the right hand side of the previous expression is positive if $\tilde{x}$ is not s-optimal. Thus, making α sufficiently large (note that the magnitude of α is not restricted in the above argument if ε is chosen small enough) the whole previous expression can be made positive. This contradicts the premise that choosing $\tilde{x}$ indefinitely from $t = 1$ represents an optimal intertemporal strategy. The proof of the last part of of Theorem 2 is thus complete. ■

Proof of Theorem 3:

It is first proven that if η is q-rational, there must exist a preference relation $\mathbf{R}$ on $\mathbb{X}$ that rationalizes η and whose asymmetric part $\mathbf{P}$ is acyclic and weakly representable (c.f. Footnote 5). If η is a q-rational pattern of behavior, there must exist some VNM utility function $V(\cdot)$, diversity function $d(\cdot)$, and perceptions $\{F(x, d)\}_{(x, d) \in \mathbb{X} \times \mathbb{R}_+}$ such that, for all $\theta \in \Theta$:

$$\eta(\theta) = \left\{ x \in \Gamma(\theta): \int V(\cdot) dF(x, 0) \geq \int V(\cdot) dF(y, d(x, y)) \text{ for all } y \in \Gamma(\theta) \right\}. \tag{49}$$

Define $\mathbf{R}$ as follows:

$$\forall x, y \in \mathbb{X}, x \mathbf{R} y \Leftrightarrow \int V(\cdot) dF(x, 0) \geq \int V(\cdot) dF(y, d(x, y)). \tag{50}$$

The binary relation $\mathbf{R}$ is complete and obviously rationalizes η , i.e., for all $\theta \in \Theta$, $\eta(\theta) = \{x \in \Gamma(\theta): x \mathbf{R} y \text{ for all } y \in \Gamma(\theta)\}$. To see that its asymmetric part $\mathbf{P}$, defined by:

$$\forall x, y \in \mathbb{X}, xPy \Leftrightarrow xRy \ \& \ \neg(yRx) \quad (51)$$

is acyclic, it is enough to show that it is weakly representable. Note that:

$$xPy \Rightarrow \int V(\cdot) dF(x,0) = V(\psi(x)) > \int V(\cdot) dF(y,0) = V(\psi(y)). \quad (52)$$

Thus, the function $h: \mathbb{X} \longrightarrow \mathbb{R}$ defined by $h(x) = V(\psi(x))$ for each $x \in \mathbb{X}$ is a weak representation of P , as desired.

I now turn to the sufficiency part of the theorem. let $h: \mathbb{X} \longrightarrow \mathbb{R}$ denote a weak representation of P . Assume, without loss of generality, that $h(\mathbb{X}) \subseteq [a,1] \subset \Omega = [0,1]$, $a > 0$, and identify the function $h(\cdot)$ with the underlying outcome function $\psi(\cdot)$. Choose then a function $V: [0,1] \longrightarrow \mathbb{R}$ satisfying (RA), and a family of perceptions $\{F(x,d)\}_{(x,d) \in \mathbb{X} \times \mathbb{R}_+}$ such that:

- (i) $\forall x \in \mathbb{X}$, $F(x,0)$ satisfies (5),
- (ii) $\int V(\cdot) dF(x,1) < \min_{y \in \mathbb{X}} h(y)$.

Such family of perceptions can always be defined if $V(0)$ has been chosen sufficiently small. As a second step in the construction, choose the diversity function $d(\cdot)$ to satisfy:

$$\forall x, y \in \mathbb{X} \left[h(x) > h(y) \ \& \ yRx \right] \Rightarrow d(x,y) \geq 1. \quad (53)$$

As it may be immediately verified, it then follows that:

$$\begin{aligned} & \left\{ x \in \Gamma(\theta): \forall y \in \Gamma(\theta), xRy \right\} = \\ & = \left\{ x \in \Gamma(\theta): \forall y \in \Gamma(\theta), \int V(\cdot) dF(x,0) \geq \int V(\cdot) dF(y,d(x,y)) \right\}, \end{aligned} \quad (54)$$

for all $\theta \in \Theta$, as desired. ■

EN BLANCO

REFERENCES

Aghion, P., P. Bolton, C. Harris & B. Jullien [1991]: "Optimal learning by experimentation", Review of Economic Studies, 58, 621-54.

Bewley, T. [1986,87]: "Knightian Decision Theory, Parts I & II", Cowles Discussion Paper no. 807.

Fishburn, P.C. [1970]: "Suborders on Commodity Spaces", Journal of Economic Theory, 2, 321-328.

Fishburn, P.C. & R.G. Vickson [1978]: "Theoretical foundations of stochastic dominance", in Whitmore, G.A., Findlay, M.C. (eds.), Stochastic Dominance, Heath: Lexington, Mass.

Grossman, S.J., R.E. Kihlstrom & L.J. Mirman [1977]: "A Bayesian approach to the production of information and learning by doing", Review of Economic Studies, 44, 533-47.

Houthakker, H. S. [1950]: "Revealed Preference and the Utility Function", Economica, N.S., 7, 159-174.

Kahneman, D. & A. Tversky [1979]: "Prospect Theory: an Analysis of decision Theory under Risk", Econometrica, 47, 263-92.

Karlin, S. & H.M. Taylor [1975]: A First Course in Stochastic Processes, London: Academic Press.

McLennan, A. [1984]: "Price dispersion and incomplete learning in the long run", Journal of Economic Dynamics and Control, 7, 331-47.

Moulin, H. [1985]: "Choice Functions Over a Finite Set: A Summary", Social Choice and Welfare, 2, 147-60.

Richter, M.K. [1966]: "Revealed Preference Theory", Econometrica, 34, 635-45.

Richter, M.K. [1971]: "Rational Choice" *in* Preferences, Utility and Demand, J.S. Chipman, L. Hurwicz, M.K. Richter, and H. Sonnenschein, eds., New York: Harcourt Brace Jovanovich.

Samuelson, P. [1938]: "A Note on the Pure Theory of Consumer's Behavior", Economica, 5, 61-71.

Samuelson, W. & R. Zeckhauser [1988]: "Status-Quo Bias in Decision Making", Journal of Risk and Uncertainty 1, 7-59.

Vega-Redondo, F. [1993]: "Simple and Inertial Behavior: An Optimizing Decision Model with Imprecise Perceptions", Economic Theory, forthcoming.

PUBLISHED ISSUES

FIRST PERIOD

- 1 "A Metatheorem on the Uniqueness of a Solution"
T. Fujimoto, C. Herrero. 1984.
- 2 "Comparing Solution of Equation Systems Involving Semipositive Operators"
T. Fujimoto, C. Herrero, A. Villar. February 1985.
- 3 "Static and Dynamic Implementation of Lindahl Equilibrium"
F. Vega-Redondo. December 1984.
- 4 "Efficiency and Non-linear Pricing in Nonconvex Environments with Externalities"
F. Vega-Redondo. December 1984.
- 5 "A Locally Stable Auctioneer Mechanism with Implications for the Stability of General Equilibrium Concepts"
F. Vega-Redondo. February 1985.
- 6 "Quantity Constraints as a Potential Source of Market Inestability: A General Model of Market Dynamics"
F. Vega-Redondo. March 1985.
- 7 "Increasing Returns to Scale and External Economies in Input-Output Analysis"
T. Fujimoto, A. Villar. 1985.
- 8 "Irregular Leontief-Straffa Systems and Price-Vector Behaviour"
I. Jimenez-Raneda / J.A. Silva. 1985.
- 9 "Equivalence Between Solvability and Strictly Semimonotonicity for Some Systems Involving Z-Functions"
C. Herrero, J.A. Silva. 1985.
- 10 "Equilibrium in a Non-Linear Leontief Model"
C. Herrero, A. Villar. 1985.
- 11 "Models of Unemployment, Persistent, Fair and Efficient Schemes for its Rationing"
F. Vega-Redondo. 1986.
- 12 "Non-Linear Models without the Monotonicity of Input Functions"
T. Fujimoto, A. Villar. 1986.
- 13 "The Perron-Frobenius Theorem for Set Valued Mappings"
T. Fujimoto, C. Herrero. 1986.
- 14 "The Consumption of Food in Time: Hall's Life Cycle Permanent Income Assumptions and Other Models"
F. Antoñanzas. 1986.
- 15 "General Leontief Models in Abstract Spaces"
T. Fujimoto, C. Herrero, A. Villar. 1986.

- 16 "Equivalent Conditions on Solvability for Non-Linear Leontief Systems"
J.A. Silva. 1986.
- 17 "A Weak Generalization of the Frobenius Theorem"
J.A. Silva. 1986
- 18 "On the Fair Distribution of a Cake in Presence of Externalities"
A. Villar. 1987.
- 19 "Reasonable Conjectures and the Kinked Demand Curve"
L.C. Corchón. 1987.
- 20 "A Proof of the Frobenius Theorem by Using Game Theory"
B. Subiza. 1987.
- 21 "On Distributing a Bundle of Goods Fairly"
A. Villar. 1987.
- 22 "On the Solvability of Complementarity Problems Involving Vo-Mappings and its Applications
to Some Economic Models"
C. Herrero, A. Villar. 1987.
- 23 "Semipositive Inverse Matrices"
J.E. Peris. 1987.
- 24 "Complementary Problems and Economic Analysis: Three Applications"
C. Herrero, A. Villar. 1987.
- 25 "On the Solvability of Joint-Production Leontief Models"
J.E. Peris, A. Villar. 1987.
- 26 "A Characterization of Weak-Monotone Matrices"
J.E. Peris, B. Subiza. 1988.
- 27 "Intertemporal Rules with Variable Speed of Adjustment: An Application to U.K.
Manufacturing Employment"
M. Burgess, J. Dolado. 1988.
- 28 "Orthogonality Test with De-Trended Data's Interpreting Monte Carlo Results using Nager
Expansions"
A. Banerjee, J. Dolado, J.W. Galbraith. 1988.
- 29 "On Lindhal Equilibria and Incentive Compatibility"
L.C. Corchón. 1988.
- 30 "Exploiting some Properties of Continuous Mappings: Lindahl Equilibria and Welfare
Egalitaria Allocations in Presence of Externalities"
C. Herrero, A. Villar. 1988.
- 31 "Smoothness of Policy Function in Growth Models with Recursive Preferences"
A.M. Gallego. 1990.
- 32 "On Natural Selection in Oligopolistic Markets"
L.C. Corchón. 1990.

- 33 "Consequences of the Manipulation of Lindahl Correspondence: An Example"
C. Bevfá, J.V. LLinares, V. Romero, T. Rubio. 1990.
- 34 "Egalitarian Allocations in the Presence of Consumption Externalities"
C. Herrero, A. Villar. 1990.

SECOND PERIOD

- WP-AD 90-01 "Vector Mappings with Diagonal Images"
C. Herrero, A. Villar. December 1990.
- WP-AD 90-02 "Langrangean Conditions for General Optimization Problems with Applications to Consumer Problems"
J.M. Gutierrez, C. Herrero. December 1990.
- WP-AD 90-03 "Doubly Implementing the Ratio Correspondence with a 'Natural' Mechanism"
L.C. Corchón, S. Wilkie. December 1990.
- WP-AD 90-04 "Monopoly Experimentation"
L. Samuelson, L.S. Mirman, A. Urbano. December 1990.
- WP-AD 90-05 "Monopolistic Competition : Equilibrium and Optimality"
L.C. Corchón. December 1990.
- WP-AD 91-01 "A Characterization of Acyclic Preferences on Countable Sets"
C. Herrero, B. Subiza. May 1991.
- WP-AD 91-02 "First-Best, Second-Best and Principal-Agent Problems"
J. Lopez-Cuñat, J.A. Silva. May 1991.
- WP-AD 91-03 "Market Equilibrium with Nonconvex Technologies"
A. Villar. May 1991.
- WP-AD 91-04 "A Note on Tax Evasion"
L.C. Corchón. June 1991.
- WP-AD 91-05 "Oligopolistic Competition Among Groups"
L.C. Corchón. June 1991.
- WP-AD 91-06 "Mixed Pricing in Oligopoly with Consumer Switching Costs"
A.J. Padilla. June 1991.
- WP-AD 91-07 "Duopoly Experimentation: Cournot and Bertrand Competition"
M.D. Alepuz, A. Urbano. December 1991.
- WP-AD 91-08 "Competition and Culture in the Evolution of Economic Behavior: A Simple Example"
F. Vega-Redondo. December 1991.
- WP-AD 91-09 "Fixed Price and Quality Signals"
L.C. Corchón. December 1991.
- WP-AD 91-10 "Technological Change and Market Structure: An Evolutionary Approach"
F. Vega-Redondo. December 1991.

- WP-AD 91-11 "A 'Classical' General Equilibrium Model"
A. Villar. December 1991.
- WP-AD 91-12 "Robust Implementation under Alternative Information Structures"
L.C. Corchón, I. Ortuño. December 1991.
- WP-AD 92-01 "Inspections in Models of Adverse Selection"
I. Ortuño. May 1992.
- WP-AD 92-02 "A Note on the Equal-Loss Principle for Bargaining Problems"
C. Herrero, M.C. Marco. May 1992.
- WP-AD 92-03 "Numerical Representation of Partial Orderings"
C. Herrero, B. Subiza. July 1992.
- WP-AD 92-04 "Differentiability of the Value Function in Stochastic Models"
A.M. Gallego. July 1992.
- WP-AD 92-05 "Individually Rational Equal Loss Principle for Bargaining Problems"
C. Herrero, M.C. Marco. November 1992.
- WP-AD 92-06 "On the Non-Cooperative Foundations of Cooperative Bargaining"
L.C. Corchón, K. Ritzberger. November 1992.
- WP-AD 92-07 "Maximal Elements of Non Necessarily Acyclic Binary Relations"
J.E. Peris, B. Subiza. December 1992.
- WP-AD 92-08 "Non-Bayesian Learning Under Imprecise Perceptions"
F. Vega-Redondo. December 1992.
- WP-AD 92-09 "Distribution of Income and Aggregation of Demand"
F. Marhuenda. December 1992.
- WP-AD 92-10 "Multilevel Evolution in Games"
J. Canals, F. Vega-Redondo. December 1992.