

Locally Strategy Proof Planning Procedures as Algorithms and Game Forms

KIMITOSHI SATO*
GRADUATE SCHOOL OF ECONOMICS
RIKKYO UNIVERSITY†
3-34-1, NISHI-IKEBUKURO, TOSHIMA-KU
TOKYO 171-8501, JAPAN

February 2009

ABSTRACT. This paper revisits the procedure developed by Sato(1983) which achieves *Aggregate Correct Revelation* in the sense that the sum of the Nash equilibrium strategies always coincides with the aggregate value of the correct MRSs. The procedure renamed *the Generalized MDP Procedure* can possess other desirable properties shared by continuous-time locally strategy proof planning procedures, i.e., feasibility, monotonicity and Pareto efficiency. Under myopia assumption, each player's dominant strategy in the local incentive game associated at any iteration of the procedure is proved to reveal his/her marginal rate of substitution for a public good. In connection with the *Generalized MDP Procedure*, this paper analyses the structure of the locally strategy proof procedures as algorithms and game forms. An alternative characterization theorem of locally strategy proof procedures is given by making use of the new Condition, Transfer It is shown that the exponent attached to the decision function of public good is characterized. Coalitional and Bayesian incentive compatibility are also discussed. Finally referred to are myopia, non-myopia and discreteness in planning procedures.

Key Words: aggregate correct revelation, Bayesian local strategy proof, coalition local strategy proof, Generalized MDP Procedure, local strategy proof, measure of incentives, Nonlinearized MDP Procedure, Fujigaki-Sato Procedure, Transfer Independence

JEL Classification: H41

1. INTRODUCTION

Since the appearance of Samuelson's seminal paper(1954), the prevalent view was that the free rider problem was inevitable in the provision of pure public goods: once the good was made available to one person, it was available to all. This pessimistic view was shattered by the advent of the MDP Procedure. It was epoch-making. Since then a large literature

*This is a paper dedicated to the XXXXth Anniversary of the MDP Procedure. An earlier version of this paper was presented at "2008 Mathematical Economics" held at Research Institute of Mathematical Sciences, Kyoto University, November 28, 2008. The author gratefully acknowledges to Professor Toru Maruyama and the participants at the workshop for their valuable comments and suggestions which substantially improved the earlier version of the paper. This revised version is prepared for the presentation at the spring meeting of the Japanese Economic Association to be held at Kyoto University, June 2009.

†e-mail address: satokt@rikkyo.ac.jp

has accumulated that develops individually rational and incentive compatible planning procedures for optimally providing public goods.

At the 1969 meeting of the Econometric Society in Brussels, Jacques Drèze and de la Vallée Poussin, and Edmond Malinvaud independently presented tâtonnement processes for guiding and financing an efficient production of public goods. As Malinvaud noted in his paper the two approaches closely resembled each other: each attempted a dynamic presentation of the Samuelson's Condition for the optimal provision of public goods. Subsequently, Malinvaud published a further article on the subject, proposing a mixed (price-quantity) procedure. Their papers are among the most important contributions in planning theory and in public economics. They came to be termed the Malinvaud-Drèze-Poussin(hereafter, MDP) Procedure, and spawned numerous papers.¹

Initiated by these three great pioneers, this field of research made remarkable progress in the last four decades. They sowed the seeds for the subsequent developments in the theory of public goods, and initiated the successful introduction of a game theoretical approach in the planning theory of public goods. Numerous succeeding contributions generated the means of providing incentives to correctly reveal preferences for public goods. The analyses of incentives in tâtonnement procedures began in late sixties and was mathematically refined by the characterization theorems of Champsaur and Rochet(1983), which generalized the previous results of Fujigaki and Sato(1981) and (1982), as well as Laffont and Maskin(1983). Champsaur and Rochet highlighted the incentive theory in the planning context to reach the acme and culminated in their generic theorems. Most of these procedures can be characterized by the conditions, the formal definitions of which are given in Section 3: (i) Feasibility, (ii) Monotonicity, (iii) Pareto Efficiency, (iv) Local Strategy Proof, and (v) Neutrality.

Very appealing for its mathematical elegance and the direct application of the Samuelson's Condition, it received a lot of attention in the 1970s and 1980s, especially on the problem of incentives in planning procedures with public goods, but there has been very little work on it over the last twenty years, leaving some very difficult problems. This paper is a follow up on the literature on the use of processes as mechanisms for aggregating the decentralized information needed for determining an optimal quantity of public goods. This paper tries to add some results on the MDP Procedure. In addition to implementation, it is required that the equilibria of the Procedure be limit points of a given dynamic adjustment process. This paper also aims at clarifying the structure of the *locally strategy proof* planning procedures as algorithms and game forms, including the MDP Procedure. They are called locally strategy proof, if players' correct revelation for a public good is a dominant strategy in the local incentive game associated with each iteration of procedures. This property is *not* possessed by the original MDP Procedure. As algorithms, they can reach any Pareto optimum. The task of the MDP Procedure is to enable the planner or the planning board to determine an optimal amount of public goods. This paper revisits the procedure developed by Sato(1983) who advocated *Aggregate Correct Revelation* in the sense that the sum of the Nash equilibrium strategies always coincides with the aggregate value of correct preferences for public goods. I could win free and escape out of the impossibility theorem among the above five desiderata, without requiring dominance. The procedure developed by Sato(1983) is able to possess similar desirable

¹See Malinvaud(1969), (1970), (1970-1971), (1971) and (1972), and Drèze and de la Vallée Poussin(1969) and (1971). For an idea of the tâtonnement process, see Drèze(1972).

features shared by continuous-time procedures, i.e., efficiency and incentive compatibility. An alternative characterization theorem of locally strategy proof procedures is given by making use of the new Condition, Transfer Independence. It means that the transfer in decision functions of public good is independent of any strategy of players.

The continuous procedures so far presented differ from that of Champsaur, Drèze, and Henry(1977) in the sense that the step-sizes for revising a public good are variable at each iteration along the solution paths.² The continuous procedures are also different from Green and Schoumaker(1978), where global information, viz., a part of each player's indifference curve, is needed to be revealed. Only local information, i.e., marginal rates of substitution(MRSs) of any player is required to determine the trajectories of the continuous processes. It is verified that the best reply strategy for each player is to reveal his/her true MRS for the public good at each iteration of procedures, which maximizes each player's payoff in the local incentive game. Thus, some continuous procedures can achieve local strategy proof. I employ the idea of modeling agents as having myopia, which can bring desirable numerous results on incentives in continuous planning procedures.

The remainder of the paper is organized as follows. The next section outlines the general framework. Section 3 reviews the MDP Procedure, renames the Non-linearized MDP Procedure, and introduces the Generalized MDP Procedure which achieve neutrality and aggregate correct revelation. It explores players' strategic manipulability in the incentive game associated with each iteration of the procedure and presents the theorems. Section 4 analyzes the structure of the locally strategy proof planning procedures. The last section provides some final remarks.

2. THE MODEL

The simplest model incorporating the essential features of the problem proposed in this paper involves two goods, one public good and one private good, whose quantities are represented by x and y , respectively. Denote y_i as an amount of the private good allocated to the i th consumer. The economy is supposed to possess n individuals. Each consumer $i \in N = \{1, \dots, n\}$ is characterized by his/her initial endowment of a private good ω_i and his/her utility function $u_i : \mathbf{R}_+^2 \rightarrow \mathbf{R}$. The production sector is represented by the transformation function $G : \mathbf{R}_+ \rightarrow \mathbf{R}_+$, where $y = G(x)$ signifies the minimal private good quantities needed to produce the public good x . It is assumed as usual that there is no production of private good. Following assumptions and definitions are used throughout this paper.

Assumption 1. For any $i \in N$, $u_i(\cdot, \cdot)$ is strictly concave and at least twice continuously differentiable.

Assumption 2. For any $i \in N$, $\partial u_i(x, y_i)/\partial x \geq 0$, $\partial u_i(x, y_i)/\partial y_i > 0$ and $\partial u_i(x, 0)/\partial x = 0$ for any x .

²The essence of the discrete version of the MDP Procedure(CDH Procedure) can be captured in Henry and Zylberberg(1977). See, in addition, Ruys(1974), Tulkens(1978), Laffont(1982) and (1985), Mukherji(1990) and Salanié(1998) for lucid summaries of the MDP Procedure. It can be seen as a *non-tâtonnement process*, due to its feasibility, one can therefore truncate it at any time. As for a contribution to the MDP literature, see Von Dem Hagen(1991), where a differential game approach is taken.

Assumption 3. $G(x)$ is convex and twice continuously differentiable.

Let $\gamma(x) = dG(x)/dx$ denote the marginal rate of transformation which is assumed to be known to the planning center. It asks each individual i to report his/her marginal rate of substitution between the public good and the private good used as a numéraire to determine an optimal quantity of the public good.

$$\pi_i(x, y_i) = \frac{\partial u_i(x, y_i)/\partial x}{\partial u_i(x, y_i)/\partial y_i}.$$

Definition 1. An allocation z is feasible if and only if

$$z \in \mathbf{Z} = \left\{ (x, y_1, \dots, y_n) \in \mathbf{R}_+^{n+1} \mid \sum_{i \in \mathbf{N}} y_i + G(x) = \sum_{i \in \mathbf{N}} \omega_i \right\}.$$

Definition 2. An allocation z is individually rational if and only if

$$(\forall i \in \mathbf{N}) [u_i(x, y_i) \geq u_i(0, \omega_i)].$$

Definition 3. A Pareto optimum for this economy is an allocation $z^* \in \mathbf{Z}$ such that there exists no feasible allocation z with

$$(\forall i \in \mathbf{N}) [u_i(x, y_i) \geq u_i(x^*, y_i^*)]$$

$$(\exists j \in \mathbf{N}) [u_j(x, y_j) > u_j(x^*, y_j^*)].$$

These assumptions and definitions altogether give us conditions for Pareto optimality in our economy.

Lemma 1. Under Assumptions 1-3, necessary and sufficient conditions for an allocation to be Pareto optimal is

$$\sum_{i \in \mathbf{N}} \pi_i \leq \gamma \quad \text{and} \quad \left(\sum_{i \in \mathbf{N}} \pi_i - \gamma \right) x = 0.$$

These are called the Samuelson's Conditions. Furthermore, conventional mathematical notation is used throughout in the same manner as in my previous paper(1983). Hereafter all variables are assumed to be functions of time t , however, the argument t is often omitted unless confusion could arise. The analyses in the following sections bypass the possibility of boundary problem at $x(t) = 0$. This is an innocuous assumption in the single public good case, because x is always increasing. The boundary problem is treated in Sato(2003). The results below can be applied to the model with many public goods.

3. THE CLASS OF MDP PROCEDURES

3.1. A Brief Review of the MDP Procedure and Its Properties

Let us describe a generic model of our planning procedures for a public good and a private good as:

$$\begin{cases} dx/dt \equiv X(t) \\ dy_i/dt \equiv Y_i(t), \forall i \in \mathbf{N}. \end{cases}$$

The MDP Procedure is the best-known member belonging to the family of the quantity-guided procedures, in which the planning center asks individual agents their MRS's between the public good and the private numéraire. Then the center revises an allocation according to the discrepancy between the sum of the reported MRSs and the MRT. The relevant information exchanged between the center and the periphery is in the form of quantity. Besides full implementation, we require an additional property: its equilibria must be approachable via an adjustment process. Suppose a game is played repeatedly in continuous time. Call $\psi(t) = (\psi_1(t), \dots, \psi_n(t)) \in \mathbf{R}_+^n$ the strategy profile played at any iteration $t \in [0, \infty)$ of the procedure. Needless to say, ψ_i is not necessarily equal to π_i , thus, the incentive problem matters.

The MDP Procedure reads:

$$\begin{cases} X(\psi(t)) = \sum_{j \in \mathbf{N}} \psi_j(t) - \gamma(t) \\ Y_i(\psi(t)) = -\psi_i(t) X(\psi(t)) + \delta_i \left\{ \sum_{j \in \mathbf{N}} \psi_j(t) - \gamma(t) \right\} X(\psi(t)), \forall i \in \mathbf{N}. \end{cases}$$

Denote a distributional coefficient $\delta_i > 0, \forall i \in \mathbf{N}$, with $\sum_{i \in \mathbf{N}} \delta_i = 1$, determined by the planner prior to the beginning of an operation of the procedure. Its role is to share among individuals the "social surplus", $\{\sum_{j \in \mathbf{N}} \psi_j(t) - \gamma(t)\} X(\psi(t))$, which is always positive except at the equilibrium.

Remark 1. Drèze and de la Vallée Poussin(1971) set $\delta_i > 0$, which was followed by Roberts(1979a,b), whereas $\delta_i \geq 0$ was assumed by Champsaur(1976) who advocated a notion of neutrality to be explained below.

A local incentive game associated with each iteration of the process is formally defined as the normal form game $(\Psi, \mathbf{U})$; $\Psi = \times_{j \in \mathbf{N}} \Psi_j \subset \mathbf{R}_+$ is the Cartesian product of the Ψ_j , which is the set of player j 's strategies, and $\mathbf{U} = (U_1, \dots, U_n)$ is the n -tuple of payoff functions. The time derivative of consumer i 's utility is such that

$$\begin{aligned} \frac{du_i}{dt} &\equiv U_i(\psi(t)) = \frac{\partial u_i}{\partial x} X(\psi(t)) + \frac{\partial u_i}{\partial y_i} Y_i(\psi(t)) \\ &= \frac{\partial u_i}{\partial x} \{ \pi_i X(\psi(t)) + Y_i(\psi(t)) \} \end{aligned}$$

which is the payoff that each player obtains at iteration t in the local incentive game along the procedure.

The behavioral hypothesis underlying the above equations is the following *myopia assumption*. In order to maximize his/her instantaneous utility increment $U_i(\psi(t))$ as his/her payoff, each player determines his/her dominant strategy $\psi_i \in \Psi_i$. Let $\psi_{-i} = (\psi_1, \dots, \psi_{i-1}, \psi_{i+1}, \dots, \psi_n) \in \Psi_{-i} = \times_{j \in \mathbf{N} - \{i\}} \Psi_j$.

Definition 4. A dominant strategy for each player in the local incentive game $(\Psi, \mathbf{U})$ is the strategy $\tilde{\psi}_i \in \Psi_i$ such that

$$(\forall \psi_i \in \Psi_i) (\forall \psi_{-i} \in \Psi_{-i}) (\forall i \in \mathbf{N}) \left[u_i(\tilde{\psi}_i, \psi_{-i}) \geq u_i(\psi_i, \psi_{-i}) \right].$$

In the Procedure below, the planning authority plans to provide an optimal quantity of a public good by revising its quantity at iteration $t = [0, \infty)$. In order for the planner to decide in what direction an allocation should be changed, it proposes a tentative feasible quantity of the public good, $x(0)$ at the initial time 0 given by the planner to which agents are asked to report his/her true MRS, $\pi_i(x(0), \omega_i)$, $\forall i \in N$, as a local privately held information. At each date t the planner can easily calculate for any t the sum of their announced MRS's to change the allocation at the next iteration $t + dt$. It is supposed that the planner can get an exact value of MRT.

The continuous-time dynamics is summarized as follows.

Step 0) At initial iteration 0, the planner proposes a feasible allocation and asks individual players to reveal their preference for the public good.

Step t) At each iteration t , players reveal their strategy and the planner calculates the discrepancy between the sum of MRS's and the MRT. Unless the equality between the above two holds, the planner suggests a new proposal allocation, and players update and reveal their preferences. If the Samuelson's Condition holds at some iteration, the MDP Procedure is truncated and an optimal quantity of the public good is determined.

3.2. Normative Conditions for the Family of the MDP Procedures

The conditions presented in Introduction are in order.

Condition F. Feasibility

$$(\forall t \in [0, \infty)) \left[\gamma(t) X(\psi(t)) + \sum_{j \in N} Y_j(\psi(t)) = 0 \right].$$

Condition M. Monotonicity

$$(\forall \psi \in \Psi) (\forall i \in N) (\forall t \in [0, \infty)) \\ \left[U_i(\psi(t)) = \frac{\partial u_i}{\partial y_i} \{ \pi_i(t) X(\psi(t)) + Y_i(\psi(t)) \} \geq 0 \right].$$

Condition PE. Pareto Efficiency

$$(\forall \psi \in \Psi) \left[X(\psi(t)) = 0 \iff \sum_{j \in N} \psi_j(t) = \gamma(t) \right].$$

Condition LSP. Local Strategy Proof

$$(\forall \psi_i \in \Psi) (\forall \psi_{-i} \in \Psi_{-i}) (\forall i \in N) (\forall t \in [0, \infty)) \\ [\pi_i(t) X(\pi_i(t), \psi_{-i}(t)) + Y_i(\pi_i(t), \psi_{-i}(t)) \geq \pi_i(t) X(\psi(t)) + Y_i(\psi(t))].$$

Condition N. Neutrality

$$(\exists z^* \in P_0) (\exists \delta \in \Delta) (\forall z(\cdot) \in Z) \left[z^* = \lim_{t \rightarrow \infty} z(t, \delta) \right]$$

where P_0 is the set of individually rational Pareto optima(IRPO), Δ is the set of $\delta = (\delta_1, \dots, \delta_n)$, and $z(\cdot)$ is a solution of the procedure.

It was Champsaur(1976) who advocated the notion of neutrality for the MDP Procedure, and Cornet(1983) generalized it by omitting two restrictive assumptions imposed by Champsaur: i.e., (i) uniqueness of solution and (ii) concavity of the utility functions. Neutrality depends on the distributional coefficient vector δ . Remember that the role of δ is to attain any IRPO by redistributing the social surplus generated during the operation of the procedure: δ varies trajectories to reach every IRPO. In other words, the planning center can guide an allocation via the choice of δ , however, it cannot predetermine a final allocation to be achieved. This is a very important property for the non-cooperative games, since the equity considerations among players matter.³

Remark 2. Conditions except *PE* must be fulfilled for any $t \in [0, \infty)$. *PE* is based on the announced values, $\psi_i, \forall i \in \mathbf{N}$, which implies that a Pareto optimum reached is not necessarily equal to the one achieved under the truthful revelation of preferences for the public good. Condition *LSP* signifies that the truth-telling is a dominant strategy. Condition *N* means that for every efficient point $z^* \in \mathbf{Z}$ and for any initial point $z_0 \in \mathbf{Z}$, there exists δ and $z(t, \delta)$, a trajectory starting from z_0 , such that $z^* = z(\infty, \delta)$.

The MDP Procedure enjoys feasibility, monotonicity, stability, neutrality, and incentive properties pertaining to minimax and Nash strategies, as was proved by Drèze and de la Vallée Poussin(1971), and Roberts(1979a,b). The MDP Procedure as an algorithm evolves in the allocation space and stops when the Samuelson's Conditions are met so that the public good quantity is optimal, and simultaneously the private good is allocated in a Pareto optimal way: i.e., $(x^*, y_1^*, \dots, y_n^*)$ is Pareto optimal.

3.3. The Locally Strategy Proof MDP Procedure

In our context, as a planner's most important task is to achieve an optimal allocation of the public good, he or she has to collect the relevant information from the periphery so as to meet the conditions presented above. Fortunately, the necessary information is available if the procedure is locally strategy proof. It was already shown by Fujigaki and Sato(1982), however, that the incentive compatible n -person MDP Procedure cannot preserve neutrality, since $\delta_i, \forall i \in \mathbf{N}$, was concluded to be fixed, i.e., $1/n$ to accomplish LSP, keeping the other conditions fulfilled. This is a sharp contrasting result, since the class of Groves mechanisms is neutral.[See Green and Laffont(1979, pp. 75-76.)]

Fujigaki and Sato(1981) presented the *Locally Strategy Proof MDP Procedure* which reads:

$$\begin{cases} X(\psi(t)) = \left(\sum_{j \in \mathbf{N}} \psi_j(t) - \gamma(t) \right) \left| \sum_{j \in \mathbf{N}} \psi_j(t) - \gamma(t) \right|^{n-2} \\ Y_i(\psi(t)) = -\psi_i(t) X(\psi(t)) + (1/n) \left(\sum_{j \in \mathbf{N}} \psi_j(t) - \gamma(t) \right) X(\psi(t)), \forall i \in \mathbf{N}. \end{cases}$$

Remark 3. We termed our procedure the "Generalized MDP Procedure" in our paper(1981). Certainly, the public good decision function was generalized to include that of the MDP Procedure, whereas, the distributional vector was fixed to the above specific value. Thus, in order to be more precise, let me call hereafter the above procedure the

³For the concepts of neutrality associated with planning procedures, see Cornet(1977a, b, c, d) and (1979), Cornet and Lasry(1977), Rochet(1982), Sato(1983), (2003) and (2005). See also d'Aspremont and Drèze(1979) for a version of neutrality which is valid for the generic context.

Fujigaki-Sato(FS) Procedure or the *Non-linearized MDP Procedure* as contrasted with the original MDP Procedure which has a linear adjustment speed of public good. The *genuine* Generalized MDP Procedure is presented below.

The FS Procedure for optimally providing the public good has the following properties:

- i) The Procedure monotonically converges to an individually rational Pareto optimum, even if agents do not report their true valuation, i.e., MRS for the public good.
- ii) Revealing his/her true MRS is always a dominant strategy for each myopically behaving agent.
- iii) The Procedure generates in the feasible allocation space similar trajectories as the MDP Procedure with uniform distribution of the instantaneous surplus occurred at each iteration, which leaves no influence of the planning authority on the final plan. Hence, the Procedure is non-neutral.

Remark 4. The property ii) is an important one that cannot be enjoyed by the original MDP Procedure except when there are only two agents with the equal surplus share, i.e., $\delta_i = 1/2, i = 1, 2$. The result on non-neutrality in iii) can be modified by designing the Generalized MDP Procedure below. See Roberts(1979a, b) for these properties.

Theorems are enumerated without proofs which were given in Fujigaki and Sato(1981).

Theorem 1. *The FS Procedure fulfills Conditions F, M, PE and LSP. However, it cannot satisfy Condition N.*

Theorem 2. *For the FS Procedure and for any $z_0 \in \mathbf{Z}$, there exists a unique solution $z(\cdot) : [0, \infty) \rightarrow \mathbf{Z}$, which is such that $\lim_{t \rightarrow \infty} z(t)$ exists and is a Pareto optimum.*

Remark 5. For the existence of solutions to the equations with the discontinuous right-hand side, see Henry(1972) and Champsaur et al.(1977) who reproduced Castaing and Valadier(1969) and Attouch and Damlamian(1972).

3.4. Best Reply Strategy and the Nash Equilibrium Strategy

In the local incentive game the planner is assumed to know the true information of individuals, since the FS Procedure induces them to elicit it. Its operation does not even require truthfulness of each player to be a Nash equilibrium strategy, but it needs only aggregate correct revelation to be a Nash equilibrium, as was verified in Sato(1983). It is easily seen from the above discussion that the FS Procedure is not neutral at all, which means that local strategy proof impedes the attainment of neutrality. Hence, Sato(1983) proposed another version of neutrality, and Condition *Aggregate Correct Revelation(ACR)* which is much weaker than *LSP*. In order to introduce Condition ACR, I need ϕ_i as a best reply strategy given by

$$\phi_i(t) = \frac{1}{n(\delta_i - 1)} \left[(1 - n)\pi_i(t) - (1 - n\delta_i) \left(\sum_{j \neq i} \psi_j - \gamma \right) \right], \quad \forall i \in \mathbf{N}.$$

Let $a' = (a_1, \dots, a_n)$ and $a_i = (1 - n\delta_i)/(n - 1)$, then one observes

$$\left(\begin{bmatrix} 1 & \dots & 0 \\ & \dots & \\ 0 & \dots & 0 \\ & \dots & \\ 0 & \dots & 1 \end{bmatrix} + \begin{bmatrix} a_1 & \dots & a_1 \\ & \dots & \\ a_i & \dots & a_i \\ & \dots & \\ a_n & \dots & a_n \end{bmatrix} \right) \begin{bmatrix} \psi_1 \\ \vdots \\ \psi_i \\ \vdots \\ \psi_n \end{bmatrix} = \begin{bmatrix} \pi_1 \\ \vdots \\ \pi_i \\ \vdots \\ \pi_n \end{bmatrix} + \gamma \begin{bmatrix} a_1 \\ \vdots \\ a_i \\ \vdots \\ a_n \end{bmatrix}.$$

Let us solve a system of n linear equations to get a Nash equilibrium strategy. First of all, the inverse matrix is computed as:

$$(I + A)^{-1} = (I - A)/(1 + \sum_{i \in \mathbf{N}} a_i) = I - A.$$

The Nash equilibrium strategy reads

$$\begin{aligned} \Phi &= (I + A)^{-1}(\pi + a\gamma) = (I - A)(\pi + a\gamma) \\ &= \pi + a\gamma - \left(\sum_{j \in \mathbf{N}} \pi_j + \gamma \sum_{j \in \mathbf{N}} a_j \right) a \\ &= \pi - \left(\sum_{j \in \mathbf{N}} \pi_j - \gamma \right) a. \end{aligned}$$

Hence, the Nash equilibrium strategy for player i is

$$\phi_i = \pi_i - \frac{1 - n\delta_i}{n - 1} \left(\sum_{j \in \mathbf{N}} \pi_j - \gamma \right).$$

It is easily seen that

$$\phi_i = \pi_i \text{ if } \delta_i = 1/n$$

which is a requirement of LSP procedures.

3.5. Aggregate Correct Revelation and the Generalized MDP Procedures

Let $\pi = (\pi_1, \dots, \pi_n)$ be a vector of MRS's for the public good and Π be its set. Sato(1983) proposed the following:

Condition ACR. Aggregate Correct Revelation:

$$(\forall \pi \in \Pi) (\forall t \in [0, \infty)) \left[\sum_{i \in \mathbf{N}} \phi_i(\pi(t)) = \sum_{i \in \mathbf{N}} \pi_i(t) \right].$$

Remark 6. Condition ACR means that the sum of Nash equilibrium strategies, $\phi_i, \forall i \in \mathbf{N}$, always coincides with the aggregate value of the correct MRS's. Clearly, ACR only claims truthfulness in the aggregate.

I needed also the following two conditions. Let $\rho : \mathbf{R}_+^n \rightarrow \mathbf{R}_+^n$ be a permutation function and $T_i(\psi)$ be a transfer in private good to agent i .

Condition TA. Transfer Anonymity

$$(\forall \psi \in \Psi) (\forall i \in \mathbf{N}) (\forall t \in [0, \infty)) [T_i(\psi(t)) = T_i(\rho(\psi(t)))].$$

Remark 7. Condition TA says that the agent i 's transfer in private good is invariant under permutation of its arguments: i.e., the order of strategies does not affect the value of $T_i(\psi(t))$, $\forall i \in \mathbf{N}$. Sato(1983) proved that $T_i(\psi(t)) = T_i\left(\sum_{j \in \mathbf{N}} \psi_j(t) - \gamma(t)\right)$ which is an example of transfer rules.

Condition TN. Transfer Neutrality

$$(\forall z^* \in \mathbf{P}_0) (\exists T \in \Omega) (\exists z(\cdot) \in \mathbf{Z}) \left[z^* = \lim_{t \rightarrow \infty} z(t, T) \right]$$

where $T = (T_1, \dots, T_n)$ is a vector of transfer functions and Ω is its set.

Now I enumerate the properties of the Generalized MDP Procedures just renamed supra. Proofs are already given in Sato(1983), so omitted here.

Theorem 3. *The Generalized MDP Procedures fulfill Conditions ACR, F, M, PE, TA and TN. Conversely, any planning process satisfying these conditions is characterized to:*

$$\begin{cases} X(\psi(t)) = \left(\sum_{j \in \mathbf{N}} \psi_j(t) - \gamma(t) \right) \left| \sum_{j \in \mathbf{N}} \psi_j(t) - \gamma(t) \right|^{n-2} \\ Y_i(\psi(t)) = -\psi_i(t) X(\psi(t)) + T_i\left(\sum_{j \in \mathbf{N}} \psi_j(t) - \gamma(t)\right), \quad \forall i \in \mathbf{N}. \end{cases}$$

Theorem 4. *Revealing preferences truthfully in any Generalized MDP Procedure is a minimax strategy for any $i \in \mathbf{N}$. It is the only minimax strategy for any $i \in \mathbf{N}$, when $x > 0$.*

Theorem 5. *$\phi_i = \pi_i$ holds for any $i \in \mathbf{N}$ at the equilibrium of the Generalized MDP Procedures.*

Theorem 6. *For every individually rational Pareto optimum z^* , there exists a vector of transfers T and a trajectory $z(\cdot) : [0, \infty) \rightarrow \mathbf{Z}$ of the differential equations defining the Generalized MDP Procedures such that $u_i(z^*) = \lim_{t \rightarrow \infty} u_i(x(t), y_i(t))$, $\forall i \in \mathbf{N}$.*

Keeping the same non-linear public good decision function as derived from Condition LSP, Sato(1983) could state the above characterization theorem. In the sequel, I employ the Generalized MDP Procedure with $T_i\left(\sum_{j \in \mathbf{N}} \psi_j - \gamma\right) = \delta_i\left(\sum_{j \in \mathbf{N}} \psi_j - \gamma\right) X(\psi)$. Via the pertinent choice of $T_i(\cdot)$ we can make the family of the Generalized MDP Procedures, including the MDP Procedure and the FS Procedure as special members.

Remark 8. Champsaur and Rochet(1983) gave a systematic study on the family of planning procedures that are asymptotically efficient and locally strategy proof. Now we know that the class of the LSP procedures is large enough, which includes the Bowen Procedure, the Champsaur-Rochet Procedure, the Fujigaki-Sato Procedure, the Generalized Wicksell Procedure, and Laffont-Maskin Procedure as special members, as classified by Rochet(1982), Sato(2003) and (2005).

4. THE STRUCTURE OF LOCALLY STRATEGY PROOF PROCEDURES

4.1. The MDP Procedure vs. the Non-linearized MDP Procedure

The existence of Non-linearized MDP Procedures is assured by the integrability and differentiability of the decision functions which determine the procedures. The MDP Procedure has a linear decision function and its adjustment speed of public good is constant. Whereas, the Non-linearized MDP Procedure has a non-linear decision function which is a kind of a “turnpike”. If illustrated in the coordinates, when the Non-linearized MDP Procedure is located far from the origin, it runs nimbler, while its adjustment speed of public good reduces in the neighborhood of the origin. This structural difference of these procedures has made a sharp contrast about the strength of incentive compatibility. This difference stems from the integrability and differentiability of the decision function of public good.

With examples, I show that the difference between the MDP Procedure and Non-linearized MDP Procedure.

Theorem 7. *When $n \geq 3$, the MDP Procedure can be manipulated by players' strategic behaviors, whereas the Non-linearized MDP Procedure cannot.*

Proof. Let me show that the original MDP Procedure can be manipulated by players in the local incentive game associated with the procedure when there are three agents. Under the truthful revelation of preference, as a payoff to player i , the time derivative of utility is represented by

$$U_i = \delta_i \left(\sum_{j \in N} \pi_j - \gamma \right)^2 X \geq 0.$$

Let φ signify underreporting of preference on the part of player 3 with $\pi_3 > \psi_3$. Whereas, it is assumed that $\psi_1 = \pi_1$ and $\psi_2 = \pi_2$.

$$\frac{du_3^\varphi}{dt} = (\pi_3 - \psi_3)X + \delta_3 \left(\sum_{j \in N} \psi_j - \gamma \right) X \geq 0.$$

If $\sum_{j \in N} \psi_j - \gamma > 0$, then

$$\frac{du_3^\varphi}{dt} - \frac{du_3}{dt} = (\pi_3 - \psi_3) \left\{ (1 - \delta_3) \left(\sum_{j \in N} \psi_j - \gamma \right) - \delta_3 \left(\sum_{j \neq i} \psi_j + \pi_i - \gamma \right) \right\}.$$

Thus, player 3 may get more payoff by falsifying his/her preference for the public good unless $\delta_3 = 1/2$.

Specify their quasi-linear utility function as $u_1 = 2x + y_1$, $u_2 = 3x + y_2$ and $u_3 = 5x + y_3$. Then, $\partial u_i / \partial y_i = 1$, $i = 1, 2$, and 3 , $\pi_1 = \psi_1 = 2$ and $\pi_2 = \psi_2 = 3$. Suppose that the public good is produced as $g(x) = 3x$ and the $\gamma = 3$. Provided that individual 3 underreports his preference by announcing $\psi_3 = 1$ instead of his true MRS, $\pi_3 = 5$.

The Generalized MDP Procedure with three persons reads

$$\begin{cases} X = \left(\sum_{j=1}^3 \psi_j - \gamma \right) \left| \sum_{j=1}^3 \psi_j - \gamma \right| \\ Y_i = -\psi_i X + \frac{1}{3} \left(\sum_{j=1}^3 \psi_j - \gamma \right) X. \end{cases}$$

With the above numerical example, this Procedure yields $du_3^0/dt = 45 < 114.33 = du_3/dt$. Similarly, $du_3^\eta/dt = 81 < 114.33 = du_3/dt$, where η means “overreporting”, when he/she reports $\psi_3 = 7$ instead of his true value, 5. Consequently, free-riding individual 3 loses his/her payoff in the both cases of underreporting and overreporting. The Non-linearized MDP Procedure gives the payoff such that

$$U_i = (\pi_i - \psi_i)X + \frac{1}{3} \left(\sum_{j=1}^3 \psi_j - \gamma \right)^2 \left| \left(\sum_{j=1}^3 \psi_j - \gamma \right) \right|$$

where $\pi_i = \psi_i$ assures $U_i \geq 0$, $\forall i = 1, 2$ and 3 , thus, the Non-linearized MDP Procedure is locally strategy proof for three persons. This is not the property enjoyed by the original MDP Procedure. Q.E.D.

4.2. An Alternative Characterization Theorem and Transfer Independence

Next, let me give an alternative proof to **Theorem 2** in Fujigaki and Sato(1981) by making use of a new axiom. This is a modified version of the property introduced by Green and Laffont(1979), which means the equality of the increment of transfer in accordance with the marginal change of strategy. This is an important condition which is connected with equity.

Condition TI. Transfer Independence:

$$(\forall i, j \in \mathbf{N}) \left[\frac{\partial T_i(\psi)}{\partial \psi_i} = \frac{\partial T_j(\psi)}{\partial \psi_j} \right].$$

Then, the following characterization theorem holds.

Theorem 8. Any planning procedure that satisfies Conditions ACR and TI is characterized to:

$$\begin{cases} G(P) = a \left(\sum_{j \in \mathbf{N}} \psi_j - \gamma \right) \left| \sum_{j \in \mathbf{N}} \psi_j - \gamma \right|^{n-1}, \quad a \in \mathbf{R}_{++} \\ T_i(\psi) = \int G \left(\sum_{j \in \mathbf{N}} \psi_j - \gamma \right) d\psi_i + H_i(\psi_{-i}), \quad \forall i \in \mathbf{N} \end{cases}$$

where $H_i(\psi_{-i})$ is an arbitrary function independent of ψ_i .

Proof. Consider the process

$$\begin{cases} X = G(P) \\ Y = -\psi_i G(P) + \delta_i P G(P). \end{cases}$$

Using the decision function specified above yields the payoff to player i :

$$U_i = \frac{\partial u_i}{\partial y_i} \{ \pi_i G(P) - \psi_i G(P) + \delta_i P G(P) \}.$$

Differentiating with respect to ψ_i this gives

$$\frac{dU_i}{d\psi_i} = \frac{\partial u_i}{\partial y_i} \left\{ \pi_i \frac{dG(P)}{dP} - G(P) - \psi_i \frac{dG(P)}{dP} + \delta_i \left[G(P) + P \frac{dG(P)}{dP} \right] \right\} = 0.$$

As a reference, if Condition LSP holds, then

$$G(P) \frac{1 - \delta_i}{\delta_i} = P \frac{dG(P)}{dP}, \quad \forall i \in \mathbf{N}.$$

This equation holds only if $\delta_i = \delta_j, \forall i, j \in \mathbf{N}$. Consequently, local strategy proof of the MDP Procedure with two persons requires $\delta_i = 1/2, \forall i \in \mathbf{N}$. Hence, the MDP Procedure can possess LSP only for a two-person economy.

Instead, if Condition ACR holds,

$$G(P) = \frac{1}{n-1} P \frac{dG(P)}{dP}, \quad \forall i \in \mathbf{N}.$$

Solving for $G(P)$ yields

$$G(P) = aP^{n-1}, \quad a \in \mathbf{R}_{++}.$$

Since $G(P)$ is sign-preserving from **Lemma 4** in Fujigaki and Sato(1982), we finally get

$$G(P) = aP|P|^{n-2}, \quad a \in \mathbf{R}_{++}.$$

Next, let me show with Conditions ACR and TI that

$$T_i(\psi) = \int G \left(\sum_{j \in \mathbf{N}} \psi_j - \gamma \right) d\psi_i + H_i(\psi_{-i}), \quad \forall i \in \mathbf{N}.$$

The best reply strategy ϕ_i for player i is, given ψ_{-i}

$$\phi_i = \left\{ \frac{\partial G(P)}{\partial \psi_i} \right\}^{-1} \left\{ \pi_i \frac{\partial G(P)}{\partial \psi_i} - G(P) + \frac{\partial T_i(\psi)}{\partial \psi_i} \right\}, \quad \forall i \in \mathbf{N}$$

where all the partial derivatives are evaluated at $\psi_i = \pi_i$.

From Condition ACR

$$\sum_{i \in \mathbf{N}} \left\{ \frac{\partial G(P)}{\partial \psi_i} \right\}^{-1} \left\{ -G(P) + \frac{\partial T_i(\psi)}{\partial \psi_i} \right\} = 0.$$

Since $G(P)$ is symmetric with respect to ψ_i ,

$$\frac{\partial G(P)}{\partial \psi_i} = \frac{\partial G(P)}{\partial \psi_j} \neq 0.$$

Thus,

$$\sum_{i \in \mathbf{N}} \frac{\partial T_i(\psi)}{\partial \psi_i} = nG(P)$$

or

$$\frac{1}{n} \sum_{i \in \mathbf{N}} \frac{\partial T_i(\psi)}{\partial \psi_i} = G(P).$$

If Condition TI holds, then

$$\frac{\partial T_i(\psi)}{\partial \psi_i} = G(P).$$

Therefore the desired conclusion follows in a straightforward manner. Q.E.D.

Remark 9. In **Theorem 8**, without Condition TI, the function $T_i(\psi)$ cannot be uniquely determined, and thus,

$$\frac{1}{n} \left\{ \frac{\partial T_i(\psi)}{\partial \psi_i} \right\} = \delta_i G(P).$$

4.3. Measure of Incentives

I show that the exponent attached to the public good decision function has a close relationship to the number of individuals participating in the procedure and that this fact enables procedures to achieve local strategy proof.

Theorem 9. Any planning procedure fulfills LSP if and only if the exponent attached to the public good decision function is $\beta = n - 1$.

Proof. Consider the following adjustment function:

$$\begin{cases} X(\psi) = \left(\sum_{j \in \mathbf{N}} \psi_j - \gamma \right)^\beta \\ Y_i(\psi) = -\psi_i X(\psi) + (1/n) \left(\sum_{j \in \mathbf{N}} \psi_j - \gamma \right) X(\psi), \forall i \in \mathbf{N} \end{cases}$$

where $\beta \geq 1$ is a parameter.

Let me show that this procedure fulfills LSP if and only if $\beta = n - 2$. For this purpose, define a *measure of incentives* below. In the local incentive game associated with each iteration of the process, the payoff for each player i is given by

$$U_i(\psi) = \frac{\partial u_i}{\partial y_i} \left\{ \pi_i - \psi_i + \frac{1}{n} \left(\sum_{j \in \mathbf{N}} \psi_j - \gamma \right) \right\} \left(\sum_{j \in \mathbf{N}} \psi_j - \gamma \right)^\beta.$$

Differentiating this with respect to ψ_i gives

$$\frac{\partial U_i(\psi_i, \psi_{-i})}{\partial \psi_i} = \frac{\partial u_i}{\partial y_i} \left\{ \beta(\pi_i - \psi_i) + \frac{\beta - n + 1}{n} \right\} \left(\sum_{j \in \mathbf{N}} \psi_j - \gamma \right)^\beta = 0.$$

Since $\left(\sum_{j \in \mathbf{N}} \psi_j - \gamma \right)^\beta \neq 0$ out of equilibrium, the best reply strategy for player i is

$$\psi_i = \pi_i + \frac{\beta - n + 1}{\beta n}.$$

Here introduced is a *measure of incentives*:

$$\Phi(n) = \sum_{i \in \mathbf{N}} (\psi_i - \pi_i)^2.$$

Substitution yields

$$\Phi(n) = \left(\frac{\beta - n + 1}{\beta n} \right)^2.$$

Differentiating this with respect to ψ_i gives

$$\frac{\partial \Phi(n)}{\partial \psi_i} = \frac{2(\beta + 1)(n - 1 - \beta)}{\beta^2 n^3} = 0.$$

The measure of incentives $\Phi(n)$ has a maximum at $n - 1$. As $n > 1$, we know that $\Phi(n) \rightarrow 0$ as $\beta \rightarrow n - 1$ and that $\Phi(n) = 0$ if and only if $\beta = n - 1$. Consequently, the Non-linearized MDP Procedure has the unique form of decision function with $\beta = n - 1$ of public good to achieve LSP. Q.E.D.

4.4. Coalitionally Locally Strategy Proof Procedures

The problem of misrepresenting preferences by colluding individuals has been dealt with for static revelation mechanisms by some authors. For instance, Bennett and Conn(1977) considered an economy with one public good and proved that there is no revelation mechanism which is group incentive compatible; that is, for any revelation mechanism to provide public goods, if any coalition formation is possible, some group of individuals will be able to gain by misrepresenting their preferences for the public good. Green and Laffont(1979) also studied the problem of coalitional manipulability. They verified under the separability of utility functions that revelation of the truth was a dominant strategy for each individual in demand revealing mechanisms used to provide public goods. They also showed that any revelation mechanism can be manipulated by coalitions of two or more agents. Their payoff by colluding, however, approaches zero as the number of agents becomes infinite, i.e., the large economy.

The main purpose of this subsection is to show whether the Local Strategy Proof MDP Procedure is robust to coalitional manipulation of preferences on the part of the agents. If the structure of coalitions is fixed and known to the planner, their misreporting can be overcome by treating each coalition as an individual agent and applying the LSP MDP Procedure to the strategies composing of the aggregated preferences over the members of each coalition. Thus, we can have a *Coalitionally Locally Strategy Proof* (CLSP) planning procedure, to be defined below. But what could happen if the coalition structure is flexible and unknown to the planner? Is it possible to construct a CLSP planning process?

Retaining the same assumptions as in Sato(1983), we add some new definitions and notation. Let $\mathbf{C} \subseteq \mathbf{N}$ be a coalition of individual agents. The vector ψ_C denotes the projection of $\psi \in \mathbf{R}^n$, the marginal rate of substitution *announced* by the coalition C . Let $\Pi_C \in \mathbf{R}^n$ be a vector of the true rate of substitution of the coalition C . We use (ψ/ψ_C) to signify the components of ψ with the exception of ψ_i , $i \in C$, and we use also the notation $(\psi, \psi_{N/C})$.

Definition 5. A joint strategy for a coalition C , $\tilde{\psi}_C \in \mathbf{R}^{|C|}$ is called a *dominant joint strategy* if it fulfills

$$(\forall \tilde{\psi}_C \in \mathbf{R}^{|C|}) (\forall \psi_{N/C} \in \mathbf{R}^{|N/C|}) (\forall i \in C) [u_i(\tilde{\psi}_C, \psi_{N/C}) \geq u_i(\psi_C, \psi_{N/C})]$$

where $||$ means a cardinality.

Definition 6. The payoff function of an agent in a coalition is given by

$$\begin{aligned} u_i(\psi_C, \psi_{N/C}) &= \frac{\partial u_i}{\partial x} X(\psi_C, \psi_{N/C}) + \frac{\partial u_i}{\partial y_i} Y_i(\psi_C, \psi_{N/C}) \\ &= \frac{\partial u_i}{\partial y_i} [\pi_i X(\psi_C, \psi_{N/C}) + Y_i(\psi_C, \psi_{N/C})]. \end{aligned}$$

Definition 7. ψ_C is said to be a *coalitionally dominant equilibrium* if it composes a dominant joint strategy against every coalition $C \in 2^n - \{\phi\}$.

Thus, we can state the condition related to coalitions.

Condition CLSP: Coalitionally Local Strategy Proof

$$\begin{aligned} &(\forall \psi_C \in \mathbf{R}^{|C|}) (\forall \psi_{N/C} \in \mathbf{R}^{|N/C|}) (\forall \psi_i \in \Psi_i) (\forall \psi_{-i} \in \Psi_{-i}) (\forall i \in C) (\forall t \in [0, \infty)) \\ &[\pi_i X(\pi_C, \psi_{N/C}) + Y_i(\pi_C, \psi_{N/C}) \geq \pi_i X(\psi_C, \psi_{N/C}) + Y_i(\psi_C, \psi_{N/C})]. \end{aligned}$$

The following theorem shows the non-existence of CLSP procedures.

Theorem 10. *There exists no procedure which fulfills Condition CLSP.*

Proof. Clearly, a CLSP planning procedure is a LSP process. Let us consider the joint payoff $U_{ik}(\psi_C, \psi_{N/C})$ of the two-size coalition $\{i, k\}$.

$$U_{ik}(\psi_C, \psi_{N/C}) = \sum_{\ell=i,k} \frac{\partial u_\ell}{\partial y_\ell} \left\{ \pi_\ell - \psi_\ell + \frac{1}{n} \left(\sum_{j \in N} \psi_j - \gamma \right) \right\} X(\psi_C, \psi_{N/C}).$$

Differentiation with respect to ψ_i gives

$$\begin{aligned} \frac{\partial U_{ik}(\psi_C, \psi_{N/C})}{\partial \psi_i} &= \sum_{\ell=i,k} \frac{\partial u_\ell}{\partial y_\ell} \left\{ \frac{1-n}{n} X(\psi_C, \psi_{N/C}) \right. \\ &\quad \left. + \left(\pi_\ell - \psi_\ell + \frac{1}{n} \sum_{j \in N} \psi_j - \frac{1}{n} \gamma \right) \frac{\partial X(\psi_C, \psi_{N/C})}{\partial \psi_\ell} \right\}. \end{aligned}$$

Since $X(\psi_C, \psi_{N/C}) = 0$ at an equilibrium where the above equation is zero if

$$\pi_i - \psi_i + \frac{1}{n} \sum_{j \in N} \psi_j - \frac{1}{n} \gamma = 0$$

and

$$\pi_k - \psi_k + \frac{1}{n} \sum_{j \in N} \psi_j - \frac{1}{n} \gamma = 0.$$

Combining these two yields

$$\pi_i - \psi_i - \pi_k + \psi_k = 0$$

which does not imply the requirement of LSP:

$$\pi_i = \psi_i \text{ and } \pi_k = \psi_k.$$

Hence, even the two-size coalition $\{i, k\}$ can manipulate the LSP procedure. Q.E.D.

4.5. Bayesian Incentive Compatible Planning Procedures

Let me refer to Bayesian strategies. A Bayesian approach to incentive compatible procedures is taken, because dominant strategies often fail to exist. Given the lack of knowledge of other players' preferences, Nash equilibrium strategies are difficult to be justified unless recontracting is permitted.

Assume that individuals' types are independently distributed; the distribution functions for individual i of type $\psi_i \in [a, b]$ is $\mu_i(\psi_i)$. These distributions are common knowledge among agents. Let $\mu_i(\psi_{-i}) \equiv \prod_{j \neq i} \mu_j(\psi_j)$ be individual i 's belief over the types of other individuals. Then, we have

Conditon BLSP: Bayesian Locally Strategy Proof

$$(\forall \psi_i \in \Psi_i) (\forall \psi_{-i} \in \Psi_{-i}) (\forall i \in N) (\forall t \in [0, \infty))$$

$$\int_{\Psi_{-i}} U_i(X(\pi_i(t), \psi_{-i}(t)), Y_i(\pi_i(t), \psi_{-i}(t))) d\mu_i(\psi_{-i}) \geq \int_{\Psi_{-i}} U_i(X(\psi_i(t), Y_i(\psi_i(t))) d\mu_i(\psi_{-i}).$$

Omitting an argument t , the following theorem is presented.

Theorem 11. *A Bayesian Locally Strategy Proof Planning Procedure is characterized as:*

$$\begin{aligned} \int_{\Psi_{-i}} X_i(\psi) d\mu_i(\psi_{-i}) &= \int_{\Psi_{-i}} \left(\sum_{j \in N} \psi_j - \gamma \right) \left| \sum_{j \in N} \psi_j - \gamma \right|^{n-2} d\mu_i(\psi_{-i}) \\ \int_{\Psi_{-i}} T_i(\psi) d\mu_i(\psi_{-i}) &= \frac{1}{n} \int_{\Psi_{-i}} \left(\sum_{j \in N} \psi_j - \gamma \right)^n d\mu_i(\psi_{-i}) \\ &\quad + \frac{1}{n(n-1)} \sum_{i \neq j} \int_{\Psi_{-i}} \left(\sum_{j \in N} \psi_j - \gamma \right)^n d\mu_i(\psi_{-i}) \end{aligned}$$

Proof: A dominant strategy is a Bayesian strategy, so that the public good decision function follows the LSP procedure as above to be the form as stated in the Theorem. The player's payoff is given by

$$\begin{aligned} U_i(t) &= \int_{\Psi_{-i}} \left\{ \frac{\partial u_i}{\partial x} X(\psi) + \frac{\partial u_i}{\partial y_i} Y_i(\psi) \right\} d\mu_i(\psi_{-i}) \\ &= \int_{\Psi_{-i}} \frac{\partial u_i}{\partial y_i} \{ \pi_i X(\psi) + Y_i(\psi) \} d\mu_i(\psi_{-i}) \\ &= \int_{\Psi_{-i}} \frac{\partial u_i}{\partial y_i} \{ (\pi_i - \psi_i) X(\psi) + T_i(\psi) \} d\mu_i(\psi_{-i}). \end{aligned}$$

Differentiating with respect to ψ_i yields the payoff:

$$U_i(t) = \int_{\Psi_{-i}} \left\{ -X(\psi) + \pi_i - \psi_i + \frac{\partial T_i(\psi)}{\partial \psi_i} \right\} d\mu_i(\psi_{-i}) = 0.$$

As required by **BLSP**, $\pi_i = \psi_i$, the above equation is

$$\int_{\Psi_{-i}} \frac{\partial T_i(\psi)}{\partial \psi_i} d\mu_i(\psi_{-i}) = \int_{\Psi_{-i}} X(\psi) d\mu_i(\psi_{-i}).$$

Integrating this with respect to ψ_i gives

$$\int_{\Psi_{-i}} T_i(\psi) d\mu_i(\psi_{-i}) = \frac{1}{n} \int_{\Psi_{-i}} \left(\sum_{j \in \mathbf{N}} \psi_j - \gamma \right)^n d\mu_i(\psi_{-i}) + H_i(\psi_{-i})$$

where $H_i(\Psi_{-i})$ is an arbitrary real valued function. In order that the sum of transfers must be zero, let me set

$$H_i(\psi_{-i}) = \frac{1}{n(n-1)} \sum_{i \neq j} \int_{\Psi_{-i}} \left(\sum_{j \in \mathbf{N}} \psi_j - \gamma \right)^n d\mu_i(\psi_{-i})$$

which is stated in the Theorem. Q.E.D.

This possibility theorem contrasts with the Roberts' impossibility theorem which is a result of dropping the myopia assumption. Roberts(1987) challenged a difficult issue which is not yet fully settled: i.e., he attempted to relax both the assumptions of myopia and complete information in a simplest version of an iterative planning framework due to Champsaur, Drèze, and Henry(1977). In his procedure the agents initially imperfectly informed but gradually learn about each other to predict future behaviors of others. He discussed the Bayesian incentive compatibility of his procedure. And he gave a numerical example of a condominium as a public good, entrance of which is redecorated by its members who use the iterative process.⁴ Much remains to be done to fully analyze the Bayesian incentive compatible planning procedures.

5. DISCUSSION ON DISCRETENESS, MYOPIA AND NONMYOPIA

Here I present some comments on the discrete procedures. Incidentally, little is known about the speed of convergence of the procedures, particularly when they are formulated in the discrete versions, which is the only realistic ones from the standpoint of actual planning. The continuous version implies that the player's responses are transmitted continuously to the planner, with no computation cost or adjustment lag.⁵ However, for the simplicity of presentation, the technical advantages of the differential approach is well-known. As Malinvaud(1970-71, p.192) rightly pointed out that a continuous formulation removes the difficult question of choosing an adjustment speed. Hence, the continuous

⁴See Spagat(1995) for incisive critics on iterative planning theory and his re-examination of the standard procedures in the Bayesian learning real-time model.

⁵See Laffont and Saint-Pierre(1979) for an exception with an information processing cost.

version is justified mainly by convenience. Moreover, a continuous formulation might be considered as an approximation to a discrete representation.⁶

Casual observations suggest that discrete procedures are more realistic than continuous ones, and that revisions of resource allocation are essentially made in discrete time. But most planning procedures discussed in the literature are formulated in continuous time, because of the difficulties involved in using the discrete version. As indicated by Malinvaud(1967) and others, this dilemma concerns a traditional technical difficulty which is summarized in such a way that if one selects a pitch large enough to get a rapid convergence, one runs the risk of no convergence. On the other hand, if one chooses a pitch small enough to expect an exact convergence, there is a possibility of delay.

Discrete versions of the MDP Procedure have been presented by several authors, and there are different strains of the related literature. The first strain – taken by Champsaur, Drèze, and Henry(1977) – is characterized by a decreasing adjustment pitch(or step-size) as a parameter, with which they could overcome a dilemma associated with a discrete formulation by keeping the pitch constant as long as it allows progress in efficiency, and by halving it as soon as it is impossible. The above-mentioned dilemma associated with discrete procedures is therefore overcome.⁷ Discussions of incentives in discrete-time MDP Procedures are given in Henry(1979), and Schoumaker(1979), (1977) and (1979). They analyzed players' strategic behaviors in the discrete MDP Processes, by ruling out the assumption of truthful revelation. The result they achieved is that their procedures still converge to a Pareto optimum even under strategic preference revelation *à la* Nash.

Approaching the same issue from another angle, Green and Schoumaker(1980) presented a discrete MDP Process with a flexible step-size at each iteration, and studied its incentive properties in the game theoretical framework. Their analysis dispensed with the "strategic indifference" assumption imposed by Henry(1979) and Schoumaker(1979): i.e., the players choose truth-telling if the resulting outcome would be indifferent. Their discrete-time procedure, however, requires reporting global information with respect to the preferences of consumers. More precisely, consumers' marginal willingness to pay functions are constrained to be compatible with a part of their utility functions. Essentially, a Nash equilibrium concept is employed. Although their ideas are interesting, the informational burden in their model is much greater than that in other approaches.

Mas-colell(1980) proposed a voluntary financing process, which is a global analog of

⁶The essence of the discrete version of the MDP Procedure(CDH Procedure) can be captured in Henry and Zylberberg(1977). See, in addition, Ruys(1974) Tulkens(1978), Laffont(1982), Mukherji(1990) and Salanié(1998) for lucid summaries of the MDP Procedure. It can be seen as a "non-tâtonnement process," because of its feasibility, one can therefore truncate it at any time. As for a contribution to the MDP literature, see Von Dem Hagen(1991), where a differential game approach is taken. De Trenquale(1992) defined a dynamic mechanism different from the MDP Procedure, that implements with local dominant strategies a Pareto efficient and individually rational allocations in a general two-agent model. Chander(1993) verified the incompatibility between core convergence property and local strategy proofness. Sato(2004) designed the Hedonic MDP Procedure for optimally providing attributes which compose the goods in the new consumer theoretical context to take "quality" into consideration.

⁷See Henry and Zylberberg(1978) for graphically illustrating how the method of a decreasing pitch successfully works until a Pareto optimum is attained. Although they treated the case with increasing returns to scale, the structure is isomorphic to the model with public goods. Crémer(1983) and (1990) took another approach to treat increasing returns to scale, as well as useful ideas that can be applied for public goods. See Heal(1986) for a comprehensive account of the planning theory and the dilemma of choosing a step-size in discrete procedures. See also Henry and Zylberberg(1977) for the Heal Procedure.

the MDP Procedure.⁸ He obtained characterizations of Pareto optimal and core states in terms of valuation functions. The incentive problem was not considered. Chander(1985) presented a discrete version of the MDP Procedure and he insisted that his system is the most informationally efficient allocation mechanism, without taking any consideration on its incentive property, though. Otsuki(1978) employed the feasible direction method in the theory of discrete planning, and applied it to the MDP and the Heal Procedures by devising implementable algorithms. Again, the problem of incentives was not treated in his paper.

Allard et al.(1989) proposed definitions of temporary and intertemporal Pareto Optimality. In their paper individuals are represented by Roy-consistent expectation functions induced by their learning processes. In order to explain their concepts of expectation functions, they referred to a pure exchange MDP Process, in which the planner asks agents to evaluate present goods and to send him/her their demands. So as to value present goods, they must forecast future quantities. Thus, Allard et al.(1989) assumed that the consumers are endowed with expectation functions.

As was criticized by Coughlin and Howe(1989), none of the above discrete procedures satisfied a criterion of intertemporal Pareto optimality. Following them, only the process devised by Green and Schoumaker(1980) insinuated a possible avenue to the criterion of intertemporal Pareto optimality. Sato(2001) showed a different version of the Green and Schoumaker(1980)'s discrete process with variable step-sizes and only local informational requirement.

Incidentally, how can we justify the myopia assumption which is a crucial underpinning to obtain a lot of fruitful results in the theory of incentives, especially in the planning procedures for optimally allocating public goods? Indeed in reality people seems to be considered to behave myopically rather than farsightedly. Matthews(1982, p. 638) wrote that "myopia may be regarded as a tractable approximation, a result of "bounded rationality"." Laffont(1985, pp. 19-20) justified myopia as follows: the participants in a planning procedure always believe that it is the last step of the procedure or that they will not enter the complexities of strategic behavior for a longer time horizon. In the MDP Procedure correct revelation of preferences is a maximin strategy in the global game, as was pointed out by Drèze. As the procedure is monotone in utility functions, the worst that could happen is the termination of the procedure: in other words, the global game reduces to the local game, in which the maximin strategy consists of correctly revealing preferences. Conversely, choosing a myopic strategy reduces to adopting a maximin approach to the global game. It would be logical, however, to adopt a maximin strategy in the local game, too.

Let me introduce two justifications of myopic by Moulin(1984, pp. 131-132). The first one is to consider an isolated player who finds himself/herself so small that his/her proper choice of strategies influences the others' choice in a negligible way. The other, which completes the first, is complete ignorance where no player knows his/her opponents' utility functions; a player knows that he/she is unable to predict in what direction the change occurs. The method of Truchon(1984) to examine a nonmyopic incentive game, where each agent's payoff is a utility at the final allocation. Different from the others, Truchon introduced a "threshold" into his model to analyze agents' strategic behavior.

⁸For another global analog, see also Dubins' mechanism which is a speed transform of the MDP Procedure explained in Green and Laffont(1979).

T. Sato(1983) also investigated how the MDP Procedure works when players with individual expectation functions nonmyopically play a sequential game, by letting them forecast what allocation would be proposed over the period when he/she takes a certain path of strategies. Assume also that the agents have rational expectations on the time interval, although the latters are bounded; they not only have complete knowledge as to the planning rules of the procedure defined below, but also can at least predict an allocation to be attained at the beginning of the next interval. Champsaur and Laroque(1982, p.326)wrote that “[s]uch a situation of limited intertemporal consistency is similar to the discrete procedures.” Champsaur and Laroque(1981) and (1982) took into consideration the effects of the agents’ strategies upon the final allocation. Sato(2001) extended his model to involve a public good in order to examine nonmyopic behaviors on the part of strategic players, as in Champsaur and Laroque(1981).

The Generalized MDP Procedure is able to keep neutrality, which is different from Champsaur and Laroque(1981)’s result on nonneutrality of the procedures with intertemporal strategic behaviors of agents. This possibility stems from Sato(1983) who proposed aggregate correct revelation as a condition to be replaceable with local strategy proofness, and he constructed a planning procedure which simultaneously satisfies three desiderata: efficiency, neutrality and aggregate correct revelation. Sato(2001) attempted a different approach, in which discussions can be extended to a piecewise linearized procedure. The above dynamic system can be generalized to involve many public goods, amounts of which can be simultaneously adjusted at each iteration. This result differs from Champsaur, Drèze, and Henry(1977), in which the quantity of only one public good can be revised at each discrete date. To examine incentive properties of the procedure, an assumption of truthful revelation of preferences is omitted. Each player’s announcement, ψ_i , is not necessarily equal to his/her true MRS, π_i . Thus, π_i must have been replaced with ψ_i in the dynamic system of the λ MDP Procedure. The nonmyopia assumption is introduced for our procedure, since a discrete time framework is a weaker representation of myopia. The procedure and the game are repeated for each interval in our framework.

What associated with the above process instead of intertemporal game used by Champsaur and Laroque(1981) is so to speak a “bounded” or “piecewise” intertemporal game, since I divide the time interval in the model. A piecewise intertemporal game played at discrete dates of each time interval of the procedure is formally defined as the normal form game $(\Psi, \mathbf{V})$. $\Psi = \times_{i \in \mathbf{N}} \Psi_i \subset \mathbf{R}_+$ is the Cartesian product of Ψ_i which is the set of player i ’s strategies, and $\mathbf{V} = (V_1(\tau_{s+1}), \dots, V_n(\tau_{s+1}))$ is the n -tuple of payoff functions at the end of the current time interval $[\tau_s, \tau_{s+1})$ such that $V_i(\tau_{s+1}) = u_i(x(\tau_{s+1}), y_i(\tau_{s+1}))$, $\forall i \in \mathbf{N}$. Let $\chi(t)$ and $v_i(t)$ be revisions at discrete date t of the public good and the private good, respectively.

The maximization problem for any player is as follows: $\forall \tau_{s+1} \in \mathbf{T}$ and $\forall t \in [\tau_s, \tau_{s+1})$

$$\begin{aligned} & \text{Max } V_i(\tau_{s+1}) \\ \text{s.t. } & x(t) = x(\tau_s) + \chi(t) \quad \text{and} \quad y_i(t) = y_i(\tau_s) + v_i(t). \end{aligned}$$

The behavioral hypothesis underlying the above equation is the nonmyopia assumption: i.e., each player determines his/her best reply strategy at the beginning of each interval $[\tau_s, \tau_{s+1})$ in order to maximize his/her payoff, $V_i(\tau_{s+1})$, at the beginning of the next interval $[\tau_{s+1}, \tau_{s+2})$.

Another Myopia Assumption: Every player is assumed to behave nonmyopically: viz., when each player determines his/her strategy in a piecewise intertemporal game, he/she does not maximize the time derivative of utility function but the utility increment based on the allocation that he/she can foresee to get at the end of the current interval.

This behavioral hypothesis may be justified by considering that the future development of an allocation cannot be predicted for exactly. Hence, every player has to make a piecewise decision under uncertainty. Players are rather assumed to forecast at least what will happen at the next discrete date. The myopia assumption is common in local games associated with both continuous and discrete planning procedures such as the MDP and the CDH(Champsaur-Drèze-Henry) Procedures. See Henry(1979), Schoumaker(1977) and (1979) for the details of this point. Also, nontâtonnement procedures are of concern in real economic life. Hence, in view of obvious practical relevance, Sato(2001) constructed our discrete process in a nontâtonnement setting, however, I was confined myself to develop a piecewise linearized process as an approximation. Under nonmyopia assumption, sincere revelation of preference for the public good at any discrete date of the Generalized MDP Procedure is a best reply strategy for each player.

6. FINAL REMARKS

The present paper has revisited the Generalized MDP Procedures and analyzed their properties. In doing so, I have extended the Sato's(1983) Procedure with a public good. In the local game associated with any iteration of the procedure, each player's payoff is the utility increment at each point of time. Laffont's differential method is used to formalize the procedure that has desirable properties. Calling this process the Nonlinearized MDP Procedure or Fujigaki-Sato Procedure, I have shown that it can simultaneously achieve efficiency and local strategy proofness. That is, it converges to a Pareto optimum and that the best replay strategy of each player at each iteration is to declare his/her true MRS, i.e., $\tilde{\psi}_i(t) = \pi_i(t)$. Instead, the Generalized MDP Procedure can possess aggregate correct revelation.

Recognizing the difficulties concerning the possibility of manipulating private information by individuals, the literature has verified that this incentive problem could be treated by the planning procedures that require a continuous revelation of information, provided that agents adopt a myopic behavior. Whereas, if individuals are farsighted, the traditional impossibility results occur, i.e., incentive compatibility is incompatible with efficiency, as were pointed out by Champsaur, Laroque and Rochet. This paper has studied an instantaneous situation where agents are only asked to reveal their true MRS at continuous dates, where the direction and speed of adjustment are changed. Consequently, the associated dynamic process named the Fujigaki-Sato Procedure has concluded to be nonlinearized. Individuals are assumed to take myopic behaviors at each date. Their behavior is hence characterized myopia, not farsightedness. The idea of looking at an intermediate time horizon for agents' manipulations of information is more natural and more realistic, but more difficult than myopia and farsightedness.

In the literature on the problem of incentives in planning procedures, the myopic strategic behavior prevailed. Many papers imposed this behavioral hypothesis; i.e., myopia, on which the forgoing discussions crucially depended, spawning numerous desirable

results in connection with the family of MDP Procedures. The aim of this paper has been to examine the consequences of the assumption that individuals choose their strategies to maximize an instantaneous change in utility function at each iteration along the procedure, as analyzed by Sato(1983). Also verified is that the Generalized MDP Procedure can always keep neutrality which is different from Champsaur and Laroque(1981) and (1982), and Laroque and Rochet(1983). They analyzed the properties of the MDP Procedure under the nonmyopic assumption. They treated the case where each individual attempts to forecast the influence of his/her announcement to the planning center over a predetermined time horizon, and optimizes his/her responses accordingly. It is proved that, if the time horizon is long enough, any noncooperative equilibrium of intertemporal game attains an approximately Pareto optimal allocation. But at such an equilibrium, the influence of the center on the final allocation is negligible, which entails nonneutrality of the procedure. Their attempt is to bridge the gap between the local instantaneous game and the global game, as was pointed out by Hammond(1979). Sato(2001) aimed, however, to bridge the gap between the local game and intertemporal game, by constructing a compromise of continuous and discrete procedures: i.e., the piecewise linearized procedure.

Acknowledgement

The author thanks Jacques Drèze, Henry Tulkens, Claude Henry and Jean-Jacques Laffont for their encouragement in my research. Laffont's too early passing is extremely lamentable.

REFERENCES

- [1] ATTOUCH, H. and A. DAMLAMIAN(1972), "On Multivalued Evolution Equations in Hilbert Spaces", *Israel Journal of Mathematics*, **12**, 373-390.
- [2] BENETT, E. and D. CONN(1977), "The Group Incentive Properties of Mechanisms for the Provision of Public Goods", *Public Choice*, **29**, 95-102.
- [3] CASTAING, C. and M. VALADIER(1969), "Equations Différentielles Multivoques dans les Espaces Localement Convexes", *Revue Française d'Informatique et de Recherche Opérationnelle*, **16**, 3-16.
- [4] CHAMPSAUR, P.(1976), "Neutrality of Planning Procedures in an Economy with Public Goods", *Review of Public Economics*, **43**, 293-300.
- [5] CHAMPSAUR, P., J. DRÈZE, and C. HENRY(1977), "Stability Theorems with Economic Applications", *Econometrica*, **45**, 272-294.
- [6] CHAMPSAUR, P. and G. LAROQUE(1981), "Le Plan Face aux Comportements Stratégiques des Unités Décentralisées", *Annales de l'INSEE*, **42**, 19-33.
- [7] CHAMPSAUR, P. and G. LAROQUE(1982), "Strategic Behavior and Decentralized Planning Procedures", *Econometrica*, **50**, 325-344.
- [8] CHAMPSAUR, P. and J.-C. ROCHET(1983), "On Planning Procedures which are Locally Strategy Proof", *Journal of Economic Theory*, **30**, 353-369.

- [9] CORNET, B.(1977a), "An Abstract Theorem for Planning Procedures", AUSLENDER, A.(ed.), *Convex Analysis and Its Application*, Lecture Notes in Economics and Mathematical Systems, **144**, Springer-Verlag, 53-59.
- [10] CORNET, B.(1977b), "Accessibilités des Optimums de Pareto par des Processus Monotones", *Comptes Rendus de l'Académie des Sciences*, **282**, Série A, 641-644.
- [11] CORNET, B.(1977c), "On Planning Procedures Defined by Multivalued Differential Equations", in *Systèmes Dynamiques et Modèles Economiques*, Colloques Internationaux du C.N.R.S., Paris, N° **259**, Chapter 2.
- [12] CORNET, B.(1977d), "Sur la Neutralité d'une Procédure de Planification", *Cahiers du Séminaire d'Économétrie*, **19**, C.N.R.S., Paris, 71-81.
- [13] CORNET, B.(1979), "Monotone Planning Procedures and Accessibility of Pareto Optima", in Aoki, M. and A. Marzollo(eds.), *New Trends in Dynamic System Theory and Economics*, Academic Press, 337-349.
- [14] CORNET, B.(1983), "Neutrality of Planning Procedures", *Journal of Mathematical Economics*, **11**, 141-160.
- [15] CORNET, B. and J.-M. LASRY(1976), "Un Théorème de Surjectivité pour une Procédure de Planification", *Comptes Rendus de l'Académie des Sciences*, **282**, Série A, 1375-1378.
- [16] D'ASPREMONT, C. and J. DRÈZE(1979), "On the Stability of Dynamic Processes in Economic Theory", *Econometrica*, **47**, 733-737.
- [17] DRÈZE, J.(1972), "A Tâtonnement Process for Investment Under Uncertainty in Private Ownership Economies", in SZEGÖ, G. and K. SHELL(eds.), *Mathematical Methods in Investment and Finance*, North-Holland, 3-23.
- [18] DRÈZE, J. and D. DE LA VALLÉE POUSSIN(1969), "A Tâtonnement Process for Guiding and Financing an Efficient Production of Public Goods", CORE Discussion Paper No. 6922; presented at the Brussels Meeting of the Econometric Society, September 1969.
- [19] DRÈZE, J. and D. DE LA VALLÉE POUSSIN(1971), "A Tâtonnement Process for Public Goods", *Review of Economic Studies*, **38**, 133-150.
- [20] FUJIGAKI, Y. and K. SATO(1981), "Incentives in the Generalized MDP Procedure for the Provision of Public Goods", *Review of Economic Studies*, **48**, 473-485.
- [21] FUJIGAKI, Y. and K. SATO(1982), "Characterization of SIIC Continuous Planning Procedures for the Optimal Provision of Public Goods", *Economic Studies Quarterly*, **33**, 211-226.
- [22] GREEN, J. and J.-J. LAFFONT(1977), "Révélation des Préférences pour les Biens Publics: Caractérisation des Mécanismes Satisfaisants", in Malinvaud, E., *Cahiers du Séminaire d'Économétrie*, No. **19**, C.N.R.S., 83-103.

- [23] GREEN, J. and J.-J. LAFFONT(1979a), *Incentives in Public Decision-Making, Studies in Public Economics, Vol. 1*, North-Holland, Amsterdam.
- [24] GREEN, J. and F. SCHOUMAKER(1980), "Incentives in Discrete-Time MDP Processes with Flexible Step-Size", *Review of Economic Studies*, **47**, 557-565.
- [25] HAMMOND, P.(1979), "Symposium on Incentive Compatibility: Introduction", *Review of Economic Studies*, **47**, 181-184.
- [26] HEAL, G.(1986), "Planning", in ARROW, K. and M. INTRILIGATOR(eds.), *Handbook of Mathematical Economics*, Vol.III, North-Holland, Amsterdam, Chapter 27.
- [27] HENRY, C.(1972), "Differential Equations with Discontinuous Right-Hand Side for Planning Procedures", *Journal of Economic Theory*, **4**, 545-551.
- [28] HENRY, C.(1972), "Non-linear Evolution Equations in Mathematical Economics", in Balakrishnan, A.(ed.), *Techniques of Optimization*, New York and London, 175-181.
- [29] HENRY, C.(1973), "Problèmes d'Existence et de Stabilité pour des Processus Dynamique Considérés en Economie Mathématique", Laboratoire d'Econométrie de l'Ecole Polytechnique, March 1973.
- [30] HENRY, C.(1976), "Vraisemblance des Mensonges et Convergence de la Procédure M.D.P.", manuscript, 1976.
- [31] HENRY, C.(1979), "On the Free Rider Problem in the M.D.P. Procedure", *Review of Economic Studies*, **46**, 293-303.
- [32] HENRY, C. and A. ZYLBERBERG(1977), "Procédures de Planification avec Rendements Croissants ou Biens Publics", *Cahiers du Séminaire d'Économétrie*, **18**, 27-38.
- [33] LAFFONT, J.-J.(1982), *Cours de Théorie Microéconomique: Vol. 1- Fondements de l'Economie Publique*, Economica, translated by BONIN, J. and H.(1985), *Fundamentals of Public Economics*, The MIT Press, Cambridge, Massachusetts.
- [34] LAFFONT, J.-J.(1985), "Incitations dans les Procédures de Planification", *Annales de l'INSEE*, **58**, 3-37.
- [35] LAFFONT, J.-J. and E. MASKIN(1983), "A Characterization of Strongly Locally Incentive Compatible Planning Procedures with Public Goods", *Review of Economic Studies*, **50**, 171-186.
- [36] LAROQUE, G. and J.-C. ROCHET(1983), "Myopic Versus Intertemporal Manipulation in Decentralized Planning Procedures", *Review of Economic Studies*, **50**, 187-196.
- [37] LASOTA, A. and Z. OPIAL, "An Application of the Kakutani-Ky Fan Theorem in the Theory of Ordinary Differential Equations", *Bulltin de l'Academie Polonaise des Sciences*, **13**, 781-786.

- [38] MALINVAUD, E.(1967), "Decentralized Procedures of Planning", in MALINVAUD, E. and M. BACHARACH(eds.), *Activity Analysis in the Theory of Growth and Planning*, Proceedings of a Conference held in Cambridge, U. K. by the International Economic Association, Macmillan.
- [39] MALINVAUD, E.(1969), "Procédures pour la Détermination d'un Programme de Consommation Collective", paper presented at the Brussels Meeting of the Econometric Society, September 1969.
- [40] MALINVAUD, E.(1970), "The Theory of Planning for Individual and Collective Consumption", INSEE, Paris, presented at Symposium on the Problems of the National Economy Modelling, June 1970, Novosibirsk.
- [41] MALINVAUD, E.(1970-71), "Procedures for the Determination of Collective Consumption", *European Economic Review*, **2**, 187-217.
- [42] MALINVAUD, E(1971), "A Planning Approach to the Public Good Problem", *Swedish Journal of Economics*, **73**, 96-121.
- [43] MALINVAUD, E.(1972), "Prices for Individual Consumption, Quantity Indicators for Collective Consumption", *Review of Economic Studies*, **39**, 385-406.
- [44] MAS-COLELL, A.(1980), "Efficiency and Decentralization in the Pure Theory of Public Goods", *Quarterly Journal of Economics*, **XCIV**, 625-641.
- [45] MATTHEWS, S.(1982), "Local Simple Games in Public Choice Mechanisms", *International Economics Review*, **23**, 623-645.
- [46] MILLERON, J.-C.(1972), "Theory of Value with Public Goods: A Survey Article", *Journal of Economic Theory*, **5**, 419-477.
- [47] MOULIN, H.(1984), "Comportement Stratégique et Communication Conflictuelle: Le Cas Non-Coopératif", *Revue Economique*, **50**, 109-125.
- [48] MUKHERJI, A.(1990), *Walrasian and Non-Walrasian Equilibria*, Clarendon Press, Oxford.
- [49] ROBERTS, J.(1979a), "Incentives in Planning Procedures for the Provision of Public Goods", *Review of Economic Studies*, **46**, 283-292.
- [50] ROBERTS, J.(1979b), "Strategic Behavior in the MDP Procedure", in LAFFONT, J.-J.(ed.), (1979), *Aggregation and Revelation of Preferences*, *Studies in Public Economics*, Vol. 2, North-Holland, Chapter 19.
- [51] ROBERTS, J.(1987), "Incentives, Information, and Iterative Planning", in GROVES, T., R. RADNER and S. REITER(eds.), *Information, Incentives, and Economic Mechanisms: Essays in Honor of Leonid Hurwicz*, Basil Blackwell, Oxford.
- [52] ROCHET, J.-C.(1982), "Planning Procedures in an Economy with Public Goods: A Survey Article", CEREMADE, DP. No. 213, Université de Paris.

- [53] RUYS, P.(1974), *Public Goods and Decentralization: The Duality Approach in the Theory of Value*, Tilberg University Press, The Netherlands.
- [54] SALANIÉ, B.(1998), *Microéconomie: Les Défaillances du Marché*, Economica, Paris; *Microeconomics of Market Failures*, The MIT Press, Cambridge, Massachusetts, 2000.
- [55] SAMUELSON, P.(1954), "The Pure Theory of Public Expenditures", *Review of Economics and Statistics*, **36**, 387-389.
- [56] SATO, K.(1983), "On Compatibility between Neutrality and Aggregate Correct Revelation for Public Goods", *Economic Studies Quarterly*, **34**, 97-109.
- [57] SATO, K.(2001), "Nonmyopia and Incentives in the Piecewise Linearized MDP Procedures with Variable Step-Sizes", presented at the Far Eastern Meeting of the Econometric Society held at the International Conference Center Kobe, Japan, July 21, 2001; the revised version presented at the autumn meeting of the Japanese Economic Association held at Hitotsubashi University, October 8, 2001.
- [58] SATO, K.(2003), "The MDP Procedure Revisited: Is It Possible to Attain Non-Samuelsonian Pareto Optima?", *Rikkyo Economic Review*, **56**, 75-105.
- [59] SATO, K.(2004), "The Hedonic MDP Procedures for Quality Attributes: Optimality and Incentives", *Rikkyo Economic Review*, **58**, 37-72.
- [60] SATO, K.(2005), "Fairness, Neutrality, and Local Strategy Proofness in Planning Procedures with Public Goods", *Rikkyo Economic Review*, **58**, 145-168.
- [61] SCHOUMAKER, F.(1976), "Best Replay Strategies in the Champsaur-Drèze-Henry Procedure", Discussion Paper 262, Centre for Mathematical Studies in Economics and Management Science, Northwestern University.
- [62] SCHOUMAKER, F.(1977), "Révélation des Préférences et Planification: une Approche Stratégique", *Recherches Economiques de Louvain*, **43**, 245-259.
- [63] TULKENS, H.(1978), "Dynamic Processes for Public Goods: A Process-Oriented Survey", *Journal of Public Economics*, **9**, 163-201; in CHANDER, P., J. DRÈZE, C. LOVELL and J. MINTZ(eds.), *Public Goods, Environmental Externalities and Fiscal Competition, Essays by Henry Tulkens*, Springer, New York, 2006, Chapter 1.
- [64] VON DEM HAGEN, O.(1991), "Strategic Behaviour in the MDP Procedure", *European Economic Review*, **35**, 121-138.