

Voting Procedure on Stopping Games of Markov Chain

Krzysztof Szajowski¹ and Masami Yasuda²

¹ Institute of Mathematics, Technical University of Wrocław, Wrocław, Poland.
szajow@im.pwr.wroc.pl

² Department of Mathematics and Informatics, Chiba University, Chiba 263, Japan.
yasuda@math.s.chiba-u.ac.jp

Abstract. The paper deals with a p person, non-cooperative game related to the observation of a Markov chain. The players observe the process up to a random moment defined by a monotonic logical function based on an individual players' decision. The concept of Nash equilibrium is used. The solution of the game for finite and infinite horizon problems is derived. A simple example is presented.

Key words. STOPPING GAME, NON-COOPERATIVE GAME, MARKOV CHAIN, VOTING PROCEDURE, MAJORITY RULE

1 Introduction

This paper deals with a p person stopping game related to the observation of a Markov chain. Let $(X_n, \mathcal{F}_n, \mathbf{P}_x)$, $n = 0, 1, 2, \dots$, be a homogeneous Markov chain defined on a probability space $(\Omega, \mathcal{F}, \mathbf{P})$ with state space $(\mathbf{E}, \mathcal{B})$. The players are able to observe the Markov chain sequentially. At each moment n their knowledge is represented by $\mathcal{F}_n$. Each player has his own utility function $f_i : \mathbf{E} \rightarrow \mathfrak{R}$, $i = 1, 2, \dots, p$, and at each moment n each player declares separately his willingness to stop the observation of the process. The process ends and the payoffs are realized when a suitable subset of players agree to it. The aim of each player is to maximize their expected payoffs. In fact, the problem will be formulated as a p person non-cooperative game with the concept of Nash equilibrium[8] as the solution. On the other hand, one can say that the multilateral stopping procedure is based on sequential voting (cf [2], [4], [15] for monotone rule concept and the mathematics of voting).

The results of this paper are mainly related to work by Kurano, Yasuda, Nakagami[6] and Yasuda, Nakagami, Kurano[16]. They have investigated the multilateral version of the optimal stopping problem for independent, identically distributed p dimensional random vectors $\overline{X_n}$. The gain function

of the i -th player is X_n^i (i -th coordinate of $\overline{X_n}$). In [6] the following class of strategies is used.

1. Each player can request a stop at any stage.
2. The majority level r ($1 \leq r \leq p$) is chosen by the players at the beginning of the game.
3. During the sequential observation process, if the number of players requesting a stop is greater than or equal to the level r , the process must be stopped.

The multilateral version of the optimal stopping problem with this class of strategies for Markov processes has been considered in [13] and [14].

This class of strategies is generalized in [16] to monotone rules. Roughly speaking, a monotone rule is a p variate, non-decreasing logical function defined on $\{0, 1\}^p$. In both papers the problem is formulated as a p person, non-cooperative game with the concept of Nash point as a solution. Paper[6] generalizes the unanimity case, i.e. $p = r$ solved by Sakaguchi[10]. The motivation for the model considered is the secretary problem (see [3] for the formulation of the problem). A solution of a bivariate version of the secretary problem is given in [6]. Presman and Sonin [9] treat this problem with another set of strategies. They considered the model in which each player's decision does not affect the stopping of the process but only his reward. Sakaguchi[11] and Kadane[5] have solved a multilateral sequential decision problem in which decisions whether to stop are made by the players alternately, instead of simultaneous decision under a monotone rule.

In the next section the formal model of the problem is formulated. The set of allowable strategies is described. Section 3 is devoted to the finite horizon case. The infinite horizon problem is investigated in Section 4. Simple examples of the games are solved in Section 5.

2 Formulation of the problem

Following the formulation of the game proposed by Kurano, Yasuda and Nakagami(cf [6], [16]) for a multilateral stopping problem with independent random vectors, this multilateral stopping problem can be described in terms of the ideas used in non-cooperative game theory(see [8], [1], [7]). Let $(X_n, \mathcal{F}_n, \mathbf{P}_x)$, $n = 0, 1, 2, \dots, N$, be a homogeneous Markov chain defined on a probability space $(\Omega, \mathcal{F}, \mathbf{P})$ with state space $(\mathbf{E}, \mathcal{B})$. The horizon can be finite or infinite ($N \in \mathbf{N} \cup \{\infty\}$, $\mathbf{N}$ denotes the set of natural numbers). The players are able to observe the Markov chain sequentially. Each player has his utility function $f_i : \mathbf{E} \rightarrow \mathbb{R}$, $i = 1, 2, \dots, p$, such that $\mathbf{E}_x |f_i(X_1)| < \infty$. If process has not stopped by n , then each player, based on $\mathcal{F}_n$, can declare independently his willingness to stop the observation of the process.

Definition 2.1 (see [16]) *An individual stopping strategy of the player i (ISS) is the sequence of random variables $\{\sigma_n^i\}_{n=1}^N$, where $\sigma_n^i : \Omega \rightarrow \{0, 1\}$, such that σ_n^i is $\mathcal{F}_n$ -measurable.*

The interpretation of the strategy is the following. If $\sigma_n^i = 1$ then player i declares that he would like to stop the process and accept the realization of X_n . Denote $\sigma^i = (\sigma_1^i, \sigma_2^i, \dots, \sigma_N^i)$ and let $\mathcal{S}^i$ be the set of ISSs of player i , $i = 1, 2, \dots, p$. Define

$$\mathcal{S} = \mathcal{S}^1 \times \mathcal{S}^2 \times \dots \times \mathcal{S}^p.$$

The element $\sigma = (\sigma^1, \sigma^2, \dots, \sigma^p)^T \in \mathcal{S}$ will be called the stopping strategy (SS). The stopping strategy $\sigma \in \mathcal{S}$ is a random matrix. The rows of the matrix are the ISSs. The columns are decisions of the players at successive moments. The actual stopping of the observation process and realization of the payoffs is defined by the stopping strategy exploiting some p -variate logical function. Let $\pi : \{0, 1\}^p \rightarrow \{0, 1\}$. The function π is called monotonic if for each $i \in \{1, 2, \dots, p\}$, $x^i \leq y^i$ implies $\pi(x^1, x^2, \dots, x^p) \leq \pi(y^1, y^2, \dots, y^p)$.

Definition 2.2 (see [16]) *A non-constant p -variate logical function π is called a monotone rule if it is a monotone function and $\pi(1, 1, \dots, 1) = 1$.*

Example 2.1 (see [16]) *The most important example of the monotone rule is the equal majority rule π_r . For $1 \leq r \leq p$ we define $\pi_r(x^1, x^2, \dots, x^p) = 1$ if $\sum_{i=1}^p x^i \geq r$ and 0 otherwise.*

In this stopping game model the stopping strategy is the list of the declaration of the individual players. The monotone rule converts the declarations to an effective stopping time.

Definition 2.3 (see [16]) *A stopping time $t_\pi(\sigma)$ generated by a SS $\sigma \in \mathcal{S}$ and a monotone rule π is defined by*

$$t_\pi(\sigma) = \inf\{1 \leq n \leq N : \pi(\sigma_n^1, \sigma_n^2, \dots, \sigma_n^p) = 1\}$$

($\inf(\emptyset) = \infty$). Since π is fixed during the analysis we skip index π and write $t(\sigma) = t_\pi(\sigma)$.

We have $\{\omega \in \Omega : t_\pi(\sigma) = n\} = \bigcap_{k=1}^{n-1} \{\omega \in \Omega : \pi(\sigma_k^1, \sigma_k^2, \dots, \sigma_k^p) = 0\} \cap \{\omega \in \Omega : \pi(\sigma_n^1, \sigma_n^2, \dots, \sigma_n^p) = 1\} \in \mathcal{F}_n$, then the random variable $t_\pi(\sigma)$ is stopping time with respect to $\{\mathcal{F}_n\}_{n=1}^N$. For any stopping time $t_\pi(\sigma)$ and $i \in \{1, 2, \dots, p\}$, let

$$f_i(X_{t_\pi(\sigma)}) = \begin{cases} f_i(X_n) & \text{if } t_\pi(\sigma) = n, \\ \limsup_{n \rightarrow \infty} f_i(X_n) & \text{if } t_\pi(\sigma) = \infty \end{cases}$$

(cf [12], [16]). If players use SS $\sigma \in \mathcal{S}$ and the individual preferences are converted to the effective stopping time by the monotone rule π , then the player i gets $f_i(X_{t_\pi(\sigma)})$.

Let $^*\sigma = (^*\sigma^1, ^*\sigma^2, \dots, ^*\sigma^p)^T$ be some fixed SS. Denote

$$^*\sigma(i) = (^*\sigma^1, \dots, ^*\sigma^{i-1}, \sigma^i, ^*\sigma^{i+1}, \dots, ^*\sigma^p)^T.$$

Definition 2.4 (see [16]) *Let the monotone rule π be fixed. The strategy $^*\sigma = (^*\sigma^1, ^*\sigma^2, \dots, ^*\sigma^p)^T \in \mathcal{S}$ is an equilibrium strategy with respect to π if for each $i \in \{1, 2, \dots, p\}$ and any $\sigma^i \in \mathcal{S}^i$ we have*

$$\mathbf{E}_x f_i(X_{t_\pi(^*\sigma)}) \geq \mathbf{E}_x f_i(X_{t_\pi(^*\sigma(i))}). \quad (2.1)$$

Remark 2.1 *In the class of games with infinite horizon we restrict ourselves to the set of stopping strategies*

$$\mathcal{S}^* = \{\sigma \in \mathcal{S} : t(\sigma) < \infty \quad \mathbf{P}_x - \text{a.e. for every } x \in \mathbf{E}\}.$$

In this class of SS the expected gain $\mathbf{E}_x f_i(X_{t(\sigma)})$ is well defined if we assume that $\mathbf{E}_x f_i^-(X_{t(\sigma)}) < \infty$, where $a^- = \max\{0, -a\}$. Hence in the infinite horizon game the class of admissible strategies is

$$\mathcal{S}_f^* = \{\sigma \in \mathcal{S}^* : \mathbf{E}_x f_i^-(X_{t(\sigma)}) < \infty \quad \text{for every } x \in \mathbf{E}, i = 1, 2, \dots, p\}.$$

To define the equilibrium strategy within the class of infinite horizon games we restrict the set of SS $\mathcal{S}$ in Definition 2.4 to the set $\mathcal{S}_f^$.*

The set of SS $\mathcal{S}$, the vector of utility functions $f = (f_1, f_2, \dots, f_p)$ and the monotone rule π define the non-cooperative game $\mathcal{G} = (\mathcal{S}, f, \pi)$. The aim of the considerations in this paper is to construct the equilibrium strategy $^*\sigma \in \mathcal{S}$ in $\mathcal{G}$. To this end we define an individual stopping set on the state space. This set describe the ISS for the player. First of all let us mention that with each ISS of the player i we can combine the sequence of stopping events $D_n^i = \{\omega : \sigma_n^i = 1\}$. For each monotone rule π there exists the corresponding set value function $\Pi : \mathcal{F} \rightarrow \mathcal{F}$ such that $\pi(\sigma_n^1, \sigma_n^2, \dots, \sigma_n^p) = \pi\{\mathbf{I}_{D_n^1}, \mathbf{I}_{D_n^2}, \dots, \mathbf{I}_{D_n^p}\} = \mathbf{I}_{\Pi(D_n^1, D_n^2, \dots, D_n^p)}$. One can say that $D_n = \{\omega : \pi(\sigma_n^1, \sigma_n^2, \dots, \sigma_n^p) = 1\} = \Pi(D_n^1, D_n^2, \dots, D_n^p)$ is the stopping event of the process at moment n . Since π is monotone, then $A^i \subseteq B^i$ implies $\Pi(A_n^1, A_n^2, \dots, A_n^p) \subseteq \Pi(B_n^1, B_n^2, \dots, B_n^p)$. For solution of this game the important class of ISS and the stopping events can be defined by subsets $C^i \in \mathcal{B}$ of the state space $\mathbf{E}$. A given set $C^i \in \mathcal{B}$ will be called the stopping set for player i at moment n if $D_n^i = \{\omega : X_n \in C^i\}$ is the stopping event.

For the logical function π we have (cf [16])

$$\pi(x^1, \dots, x^p) = x^i \cdot \pi(x^1, \dots, \overset{i}{1}, \dots, x^p) + \overline{x^i} \cdot \pi(x^1, \dots, \overset{i}{0}, \dots, x^p).$$

It implies that for $D^i \in \mathcal{F}$

$$\begin{aligned} \Pi(D^1, \dots, D^p) = & \{D^i \cap \Pi(D^1, \dots, \overset{i}{\Omega}, \dots, D^p)\} \\ & \cup \{\overline{D^i} \cap \Pi(D^1, \dots, \overset{i}{\emptyset}, \dots, D^p)\}. \end{aligned} \quad (2.2)$$

Let f_i, g_i be real valued, integrable (i.e. $\mathbf{E}_x|f_i(X_1)| < \infty$) functions defined on $\mathbf{E}$. For fixed $D_n^j, j = 1, 2, \dots, p, j \neq i$, and $C^i \in \mathcal{B}$ define

$$\psi(C^i) = \mathbf{E}_x \left[f_i(X_1) \mathbf{I}_{D_1(D_1^i)} + g_i(X_1) \mathbf{I}_{\overline{D_1(D_1^i)}} \right]$$

where ${}^iD_1(A) = \Pi(D_1^1, \dots, D_1^{i-1}, A, D_1^{i+1}, \dots, D_1^p)$ and $D_1^i = \{\omega : X_n \in C^i\}$. By (2.2) we have

$$\begin{aligned} \psi(C^i) &= \mathbf{E}_x(f_i(X_1) - g_i(X_1)) \mathbf{I}_{D_1^i \cap {}^iD_1(\Omega)} \\ &\quad + \mathbf{E}_x(f_i(X_1) - g_i(X_1)) \mathbf{I}_{\overline{D_1^i} \cap {}^iD_1(\emptyset)} + \mathbf{E}_x g_i(X_1) \\ &= \mathbf{E}_x(f_i(X_1) - g_i(X_1)) \mathbf{I}_{D_1^i \cap ({}^iD_1(\Omega) \setminus {}^iD_1(\emptyset))} \\ &\quad + \mathbf{E}_x(f_i(X_1) - g_i(X_1)) \mathbf{I}_{{}^iD_1(\emptyset)} + \mathbf{E}_x g_i(X_1) \\ &\leq \mathbf{E}_x(f_i(X_1) - g_i(X_1))^+ \mathbf{I}_{D_1^i \cap \Pi({}^iD_1(\Omega) \setminus {}^iD_1(\emptyset))} \\ &\quad + \mathbf{E}_x(f_i(X_1) - g_i(X_1)) \mathbf{I}_{{}^iD_1(\emptyset)} + \mathbf{E}_x g_i(X_1). \end{aligned}$$

Let $*C^i = \{x \in \mathbf{E} : f_i(x) \geq g_i(x)\}$ and denote $*D_n^i = \{\omega : X_n \in *C^i\}$. We get

$$\begin{aligned} \psi(*C^i) &= \mathbf{E}_x(f_i(X_1) - g_i(X_1))^+ \mathbf{I}_{*D_1(\Omega)} \\ &\quad - \mathbf{E}_x(f_i(X_1) - g_i(X_1))^- \mathbf{I}_{*D_1(\emptyset)} + \mathbf{E}_x g_i(X_1), \end{aligned}$$

where $a^+ = \max\{0, a\}$ and $a^- = \min\{0, -a\}$. This allows us to prove the following Lemma.

Lemma 2.1 *Let f_i, g_i , be integrable and let $C^j \in \mathcal{B}, j = 1, 2, \dots, p, j \neq i$, be fixed. Then the set $*C^i = \{x \in \mathbf{E} : f_i(x) - g_i(x) \geq 0\} \in \mathcal{B}$ is such that*

$$\psi(*C^i) = \sup_{C^i \in \mathcal{B}} \psi(C^i)$$

and

$$\begin{aligned} \psi(*C^i) &= \mathbf{E}_x(f_i(X_1) - g_i(X_1))^+ \mathbf{I}_{*D_1(\Omega)} \\ &\quad - \mathbf{E}_x(f_i(X_1) - g_i(X_1))^- \mathbf{I}_{*D_1(\emptyset)} + \mathbf{E}_x g_i(X_1). \end{aligned} \tag{2.3}$$

Based on Lemma 2.1 we derive the recursive formulae defining the equilibrium point and the equilibrium payoff for the finite horizon game.

3 Finite horizon game

Let the horizon N be finite. If the equilibrium strategy $*\sigma$ exists, then we denote $v_{i,N}(x) = \mathbf{E}_x f_i(X_{t(*\sigma)})$ to be the equilibrium payoff of i -th player when $X_0 = x$. For backward induction we introduce some useful notation. Let $\mathcal{S}_n^i = \{\{\sigma_k^i\}, k = n, \dots, N\}$ be the set of ISS for moments $n \leq k \leq N$

and $\mathcal{S}_n = \mathcal{S}_n^1 \times \mathcal{S}_n^2 \times \dots \times \mathcal{S}_n^p$. The SS for moments not earlier than n is ${}^n\sigma = ({}^n\sigma^1, {}^n\sigma^2, \dots, {}^n\sigma^p) \in \mathcal{S}_n$, where ${}^n\sigma^i = (\sigma_n^i, \sigma_{n+1}^i, \dots, \sigma_N^i)$. Denote

$$t_n = t_n(\sigma) = t({}^n\sigma) = \inf\{n \leq k \leq N : \pi(\sigma_k^1, \sigma_k^2, \dots, \sigma_k^p) = 1\}$$

to be the stopping time not earlier than n .

Definition 3.1 *The stopping strategy ${}^{n*}\sigma = ({}^{n*}\sigma^1, {}^{n*}\sigma^2, \dots, {}^{n*}\sigma^p)$ is an equilibrium in $\mathcal{S}_n$ if*

$$\mathbf{E}_x f_i(X_{t_n}({}^{n*}\sigma)) \geq \mathbf{E}_x f_i(X_{t_n}({}^{n*}\sigma(i))) \quad \mathbf{P}_x - a.e.$$

for every $i \in \{1, 2, \dots, p\}$, where

$${}^{n*}\sigma(i) = ({}^{n*}\sigma^1, \dots, {}^{n*}\sigma^{i-1}, {}^n\sigma^i, {}^{n*}\sigma^{i+1}, \dots, {}^{n*}\sigma^p).$$

Denote

$$v_{i,N-n+1}(X_{n-1}) = \mathbf{E}_x[f_i(X_{t_n}({}^{n*}\sigma)) | \mathcal{F}_{n-1}] = \mathbf{E}_{X_{n-1}} f_i(X_{t_n}({}^{n*}\sigma)).$$

At $n = N$ the players request a stop and $v_{i,0}(x) = f_i(x)$. Let us assume that the process is not stopped before n and the players are using the equilibrium strategies ${}^*\sigma_k^i$, $i = 1, 2, \dots, p$, at moments $k = n+1, \dots, N$. Choose player i and assume that the other players are using the equilibrium strategies ${}^*\sigma_n^j$, $j \neq i$, and player i is using strategy σ_n^i defined by some stopping set C^i . Then the expected payoff $\varphi_{N-n}(X_{n-1}, C^i)$ of player i in the game starting at moment n , when the state of the Markov chain at moment $n-1$ is X_{n-1} , is equal to

$$\varphi_{N-n}(X_{n-1}, C^i) = \mathbf{E}_{X_{n-1}} \left[f_i(X_n) \mathbf{I}_{i \in D_n(D_n^i)} + v_{i,N-n}(X_n) \mathbf{I}_{i \notin D_n(D_n^i)} \right],$$

where ${}^iD_n(A) = \Pi({}^*D_n^1, \dots, {}^*D_n^{i-1}, A, {}^*D_n^{i+1}, \dots, {}^*D_n^p)$.

By Lemma 2.1 the conditional expected gain $\varphi_{N-n}(X_{n-1}, C^i)$ attains its maximum on the stopping set ${}^*C_n^i = \{x \in \mathbf{E} : f_i(x) - v_{i,N-n}(x) \geq 0\}$ and

$$\begin{aligned} v_{i,N-n+1}(X_{n-1}) &= \mathbf{E}_x[(f_i(X_n) - v_{i,N-n}(X_n))^+ \mathbf{I}_{i \in D_n(\Omega)} | \mathcal{F}_{n-1}] \\ &\quad - \mathbf{E}_x[(f_i(X_n) - v_{i,N-n}(X_n))^- \mathbf{I}_{i \in D_n(\emptyset)} | \mathcal{F}_{n-1}] \\ &\quad + \mathbf{E}_x[v_{i,N-n}(X_n) | \mathcal{F}_{n-1}] \end{aligned} \quad (3.1)$$

$\mathbf{P}_x$ -a.e.. It allows us to formulate the following construction for the equilibrium strategy and the equilibrium value for the game $\mathcal{G}$.

Theorem 3.1 *In the game $\mathcal{G}$ with finite horizon N we have the following solution.*

- (i) *The equilibrium value $v_i(x)$, $i = 1, 2, \dots, p$, of the game $\mathcal{G}$ can be calculated recursively as follows:*

1. $v_{i,0}(x) = f_i(x)$;
2. For $n = 1, 2, \dots, N$ we have $\mathbf{P}_x$ -a.e.

$$\begin{aligned} v_{i,n}(x) = & \mathbf{E}_x[(f_i(X_{N-n+1}) - v_{i,n-1}(X_{N-n+1}))^+ \mathbf{I}_{i \cdot D_{N-n+1}(\Omega)} | \mathcal{F}_{N-n}] \\ & - \mathbf{E}_x[(f_i(X_{N-n+1}) - v_{i,n-1}(X_{N-n+1}))^- \mathbf{I}_{i \cdot D_{N-n+1}(\emptyset)} | \mathcal{F}_{N-n}] \\ & + \mathbf{E}_x[v_{i,n-1}(X_{N-n+1}) | \mathcal{F}_{N-n}], \end{aligned}$$

for $i = 1, 2, \dots, p$.

- (ii) The equilibrium strategy $^*\sigma \in \mathcal{S}$ is defined by the SS of the players $^*\sigma_n^i$, where $^*\sigma_n^i = 1$ if $X_n \in ^*C_n^i$, and $^*C_n^i = \{x \in \mathbf{E} : f_i(x) - v_{i,N-n}(x) \geq 0\}$, $n = 0, 1, \dots, N$.

We have $v_i(x) = v_{i,N}(x)$, and $\mathbf{E}_x f_i(X_{t(^*\sigma)}) = v_{i,N}(x)$, $i = 1, 2, \dots, p$.

PROOF. Part (i) is a consequence of the Lemma 2.1. We have $v_{i,1}(x) = \mathbf{E}_x f_i(X_1)$ and $v_{i,1}(X_{N-1}) = \mathbf{E}_{X_{N-1}} f_i(X_{t(^* \sigma)})$ $\mathbf{P}_x$ -a.e.. Assume that

$$v_{i,N-n-1}(X_{n+1}) = \mathbf{E}_{X_{n+1}} f_i(X_{t((n+2) \cdot \sigma)}) \quad \mathbf{P}_x - \text{a.e.} \quad (3.2)$$

and define $^*D_{n+1} = \Pi(^*D_{n+1}^1, \dots, ^*D_{n+1}^i, \dots, ^*D_{n+1}^p)$. We have

$$\begin{aligned} & \mathbf{E}_{X_n} f_i(X_{t((n+1) \cdot \sigma)}) \\ = & \mathbf{E}_{X_n} [f_i(X_{t((n+1) \cdot \sigma)}) \mathbf{I}_{\cdot D_{n+1}} + f_i(X_{t((n+1) \cdot \sigma)}) \mathbf{I}_{\overline{\cdot D_{n+1}}}] \\ = & \mathbf{E}_{X_n} f_i(X_{n+1}) \mathbf{I}_{\cdot D_{n+1}} + \mathbf{E}_{X_n} \{ \mathbf{E}_{X_{n+1}} f_i(X_{t((n+1) \cdot \sigma)}) \} \mathbf{I}_{\overline{\cdot D_{n+1}}} \\ \stackrel{(3.2)}{=} & \mathbf{E}_{X_n} f_i(X_{n+1}) \mathbf{I}_{\cdot D_{n+1}} + \mathbf{E}_{X_n} v_{i,N-n+1}(X_{n+1}) \mathbf{I}_{\overline{\cdot D_{n+1}}} \\ \stackrel{(\text{L. 2.1})}{=} & \mathbf{E}_{X_n} (f_i(X_{n+1}) - v_{i,N-n+1}(X_{n+1}))^+ \mathbf{I}_{i \cdot D_{n+1}(\Omega)} \\ & - \mathbf{E}_{X_n} (f_i(X_{n+1}) - v_{i,N-n+1}(X_{n+1}))^- \mathbf{I}_{i \cdot D_{n+1}(\emptyset)} \\ & + \mathbf{E}_{X_n} v_{i,N-n+1}(X_{n+1}). \end{aligned}$$

We show that $^*\sigma$ is an equilibrium point in the game $\mathcal{G}$. To this end let us assume, that $p-1$ players are using the strategies $^*\sigma^i$ and one of the players, say the player 1, uses the strategy $\sigma_{\{n\}}^1 = (\sigma_1^1, \sigma_2^1, \dots, \sigma_n^1, ^*\sigma_{n+1}^1, \dots, ^*\sigma_N^1)$, ($\sigma_{\{0\}}^1 = ^*\sigma^1$). We have

$$\mathbf{E}_x f_1(X_{t(\sigma_{\{n\}})}) \leq \mathbf{E}_x f_1(X_{t(\sigma_{\{n-1\}})}),$$

where $\sigma_{\{n\}} = (\sigma_{\{n\}}^1, ^*\sigma^2, \dots, ^*\sigma^p)^T$. It means that player 1 is using some fixed strategy at moments $1, 2, \dots, n$ and the strategy defined by (ii) thereafter. We get

$$\begin{aligned} \mathbf{E}_x f_1(X_{t(\sigma_{\{n\}})}) &= \mathbf{E}_x f_1(X_{t(\sigma_{\{n\}})}) \mathbf{I}_{\{t(\sigma_{\{n\}}) < n\}} \\ &\quad + \mathbf{E}_x f_1(X_{t(\sigma_{\{n\}})}) \mathbf{I}_{\{t(\sigma_{\{n\}}) \geq n\}} \\ &= \mathbf{E}_x f_1(X_{t(\sigma_{\{n-1\}})}) \mathbf{I}_{\{t(\sigma_{\{n-1\}}) < n\}} \\ &\quad + \mathbf{E}_x f_1(X_{t(\sigma_{\{n\}})}) \mathbf{I}_{\{t(\sigma_{\{n\}}) \geq n\}} \end{aligned}$$

and

$$\begin{aligned}
\mathbf{E}_x f_1(X_{t(\sigma\{n\})}) \mathbf{I}_{\{t(\sigma\{n\}) \geq n\}} &= \mathbf{E}_x f_1(X_{t_n(\bullet\sigma)}) \mathbf{I}_{\{t(\sigma\{n\}) \geq n\}} \\
&\leq \mathbf{E}_x \mathbf{E}_x [f_1(X_n) \mathbf{I}_{\bullet D_n} | \mathcal{F}_{n-1}] \mathbf{I}_{\{t(\sigma\{n\}) \geq n\}} \\
&\quad + \mathbf{E}_x \mathbf{E}_x [v_{i,N-n}(X_n) \mathbf{I}_{\overline{\bullet D_n}} | \mathcal{F}_{n-1}] \mathbf{I}_{\{t(\sigma\{n\}) \geq n\}} \\
&= \mathbf{E}_x v_{i,N-n+1}(X_{n-1}) \mathbf{I}_{\{t((n-1)\bullet\sigma) \geq n\}}.
\end{aligned}$$

Hence

$$\begin{aligned}
\mathbf{E}_x f_1(X_{t(\sigma\{n\})}) &\leq \mathbf{E}_x f_1(X_{t(\sigma\{n-1\})}) \mathbf{I}_{\{t(\sigma\{n-1\}) < n\}} \\
&\quad + \mathbf{E}_x v_{i,N-n+1}(X_{n-1}) \mathbf{I}_{\{t((n-1)\bullet\sigma) \geq n\}} \\
&= \mathbf{E}_x f_1(X_{t(\sigma\{n-1\})}) \mathbf{I}_{\{t(\sigma\{n-1\}) < n\}} \\
&\quad + \mathbf{E}_x f_1(X_{t(\sigma\{n-1\})}) \mathbf{I}_{\{t(\sigma\{n-1\}) \geq n\}} \\
&= \mathbf{E}_x f_1(X_{t(\sigma\{n-1\})}).
\end{aligned}$$

This ends the proof of the theorem. $\square$

4 Infinite horizon game

In this class of games the equilibrium strategy is defined in Definition 2.4 but only for the class of SS described in Remark 2.1. In this section a sufficient conditions for the existence of the equilibrium strategy for the infinite horizon game are formulated. Let $\bullet\sigma \in \mathcal{S}_f^*$ be an equilibrium strategy. Denote

$$v_i(x) = \mathbf{E}_x f_i(X_{t(\bullet\sigma)}).$$

Let us assume that $^{(n+1)}\bullet\sigma \in \mathcal{S}_{f,n+1}^*$ is constructed and it is an equilibrium strategy. If players $j = 1, 2, \dots, p, j \neq i$, apply at n the equilibrium strategies $\bullet\sigma_n^j$, player i the strategy σ_n^i defined by stopping set $\mathcal{C}^i$ and $^{(n+1)}\bullet\sigma$ at $n+1, n+2, \dots$, then the expected payoff of the player i , when history of the process up to moment $n-1$ is known, is given by

$$\varphi_n(X_{n-1}, C^i) = \mathbf{E}_{X_{n-1}} \left[f_i(X_n) \mathbf{I}_{\bullet D_n(D_n^i)} + v_i(X_n) \mathbf{I}_{\overline{\bullet D_n(D_n^i)}} \right],$$

where ${}^i\bullet D_n(A) = \Pi({}^i\bullet D_n^1, \dots, {}^i\bullet D_n^{i-1}, A, {}^i\bullet D_n^{i+1}, \dots, {}^i\bullet D_n^p)$, ${}^i\bullet D_n^j = \{\omega \in \Omega : \bullet\sigma_n^j = 1\}$, $j = 1, 2, \dots, p, j \neq i$, and $D_n^i = \{\omega \in \Omega : \sigma_n^i = 1\} = 1\} = \{\omega \in \Omega : X_n \in \mathcal{C}^i\}$. By Lemma 2.1 the conditional expected gain $\varphi_n(X_{n-1}, C^i)$ attains its maximum on the stopping set ${}^i\bullet C_n^i = \{x \in \mathbf{E} : f_i(x) \geq v_i(x)\}$ and

$$\begin{aligned}
\varphi_n(X_{n-1}, {}^i\bullet C^i) &= \mathbf{E}_x [(f_i(X_n) - v_i(X_n))^+ \mathbf{I}_{\bullet D_n(\Omega)} | \mathcal{F}_{n-1}] \\
&\quad - \mathbf{E}_x [(f_i(X_n) - v_i(X_n))^- \mathbf{I}_{\bullet D_n(\emptyset)} | \mathcal{F}_{n-1}] \\
&\quad + \mathbf{E}_x [v_i(X_n) | \mathcal{F}_{n-1}].
\end{aligned}$$

Let us assume that there exists solution $(w_1(x), w_2(x), \dots, w_p(x))$ of the equations

$$\begin{aligned} w_i(x) &= \mathbf{E}_x(f_i(X_1) - w_i(X_1))^+ \mathbf{I}_{i \bullet D_1(\Omega)} \\ &\quad - \mathbf{E}_x(f_i(X_1) - w_i(X_1))^- \mathbf{I}_{i \bullet D_1(\emptyset)} + \mathbf{E}_x w_i(X_1), \end{aligned} \quad (4.1)$$

$i = 1, 2, \dots, p$. Consider the stopping game with following payoff function for $i = 1, 2, \dots, p$.

$$\phi_{i,N}(x) = \begin{cases} f_i(x) & \text{if } n < N, \\ v_i(x) & \text{if } n \geq N. \end{cases}$$

Lemma 4.1 *Let $^*\sigma \in \mathcal{S}_f^*$ be an equilibrium strategy in the infinite horizon game $\mathcal{G}$. For every N we have*

$$\mathbf{E}_x \phi_{i,N}(X_{t^*}) = v_i(x).$$

PROOF. For $N = 1$ we have $\mathbf{E}_x \phi_{i,1}(X_{t^*}) = v_1(x)$ by Lemma 2.1. Our induction assumption is that for $k = 1, 2, \dots, N$

$$\mathbf{E}_x \phi_{i,k}(X_{t^*}) = v_i(x). \quad (4.2)$$

We have

$$\begin{aligned} \mathbf{E}_x \phi_{i,N+1}(X_{t^*}) &= \mathbf{E}_x \phi_{i,N+1}(X_{t^*}) \mathbf{I}_{\{t^* > 1\}} + \mathbf{E}_x \phi_{i,N+1}(X_{t^*}) \mathbf{I}_{\{t^* = 1\}} \\ &= \mathbf{E}_x [\phi_{i,N+1}(X_{t^*}) \mathbf{I}_{\bullet D_1} + \phi_{i,N+1}(X_{t^*}) \mathbf{I}_{\bullet \overline{D_1}}] \\ &= \mathbf{E}_x [f_i(X_1) \mathbf{I}_{\bullet D_1} + \mathbf{E}_{X_1} \phi_{i,N+1}(X_{t^*}) \mathbf{I}_{\bullet \overline{D_1}}] \\ &\stackrel{(4.2)}{=} \mathbf{E}_x [f_i(X_1) \mathbf{I}_{\bullet D_1} + v_i(X_1) \mathbf{I}_{\bullet \overline{D_1}}] \\ &\stackrel{(\text{L. 2.1})}{=} \mathbf{E}_x (f_i(X_1) - v_i(X_1))^+ \mathbf{I}_{i \bullet D_1(\Omega)} \\ &\quad - \mathbf{E}_x (f_i(X_1) - v_i(X_1))^- \mathbf{I}_{i \bullet D_1(\emptyset)} + \mathbf{E}_x v_i(X_1) \\ &\stackrel{(4.1)}{=} v_i(x), \end{aligned}$$

where $^*D_1 = \Pi(^*D_1^1, \dots, ^*D_1^i, \dots, ^*D_1^p)$. It ends the proof of Lemma 4.1. $\square$

Let us assume that for $i = 1, 2, \dots, p$ and every $x \in \mathbf{E}$ we have

$$\mathbf{E}_x [\sup_{n \in \mathbf{N}} f_i^+(X_n)] < \infty. \quad (4.3)$$

Theorem 4.2 *Let $(X_n, \mathcal{F}_n, \mathbf{P}_x)_{n=0}^\infty$ be homogeneous Markov chain and the payoff functions of the players fulfill (4.3). If $t^* = t(^*\sigma)$, $^*\sigma \in \mathcal{S}_f^*$ then $\mathbf{E}_x f_i(X_{t^*}) = v_i(x)$.*

PROOF. By Lemma 4.1 and by (4.3) we have

$$\begin{aligned}
& |\mathbf{E}_x f_i(X_{t^*}) - \mathbf{E}_x \phi_{i,N}(X_{t^*})| \\
&= |\mathbf{E}_x(f_i(X_{t^*}) - v_i(X_{t^*}))\mathbf{I}_{\{t^* > N\}}| \\
&\leq \mathbf{P}_x\{t^* > N\} |\mathbf{E}_x[(f_i(X_{t^*})\mathbf{I}_{\{t^* > N\}}) - v_i(X_{t^*})\mathbf{I}_{\{t^* > N\}}]| \\
&\leq \mathbf{P}_x\{t^* > N\} \{\mathbf{E}_x[\sup_{n \geq N} (f_i^+(X_{t^*})) + v_i(x)]\}.
\end{aligned}$$

To complete the proof let $N \rightarrow \infty$. $\square$

Theorem 4.3 *Let the stopping strategy $^*\sigma \in \mathcal{S}_f^*$ be defined by the stopping sets $^*C_n^i = \{x \in \mathbf{E} : f_i(x) \geq v_i(x)\}$, $i = 1, 2, \dots, p$, then $^*\sigma$ is the equilibrium strategy in the infinite stopping game $\mathcal{G}$.*

PROOF. We show that if $p-1$ players are using the strategies $^*\sigma^j$ and one player, say the i -th player, uses the strategy such as $\sigma_{\{n\}}^i = (\sigma_1^i, \sigma_2^i, \dots, \sigma_n^i, ^*\sigma_{n+1}^i, ^*\sigma_{n+2}^i, \dots)$, $n \geq 1$, $\sigma_{\{0\}}^i = ^*\sigma^i$, then

$$\mathbf{E}_x f_i(X_{t(\sigma_{\{n\}})}) \leq \mathbf{E}_x f_i(X_{t(\sigma_{\{n-1\}})}),$$

where $\sigma_{\{n\}} = (^*\sigma^1, \dots, ^*\sigma^{i-1}, \sigma_{\{n\}}^i, ^*\sigma^{i+1}, \dots, ^*\sigma^p)^T$. Denote $t(\sigma_{\{n\}}) = t(n)$. Since from stage $n-1$ the strategies $t(n)$ and $t(n-1)$ coincide, then

$$\mathbf{E}_x f_i(X_{t(n)}) = \mathbf{E}_x f_i(X_{t(n-1)})\mathbf{I}_{\{t(n-1) < n\}} + \mathbf{E}_x f_i(X_{t(n)})\mathbf{I}_{\{t(n) \geq n\}}.$$

Moreover

$$\begin{aligned}
& \mathbf{E}_x f_i(X_{t(n)})\mathbf{I}_{\{t(n) \geq n\}} \\
&= \mathbf{E}_x \mathbf{E}_x[f_i(X_{t(n)})|\mathcal{F}_{n-1}]\mathbf{I}_{\{t(n) \geq n\}} \\
&= \mathbf{E}_x \mathbf{E}_{X_{n-1}} f_i(X_{t(n)})\mathbf{I}_{\{t(n) \geq n\}} \\
&= \mathbf{E}_x \{\mathbf{E}_{X_{n-1}}[f_i(X_n)\mathbf{I}_{D_n} + f_i(X_{t(n)})\mathbf{I}_{\overline{D_n}}]\mathbf{I}_{\{t(n) \geq n\}}\} \\
&\leq \mathbf{E}_x \{\mathbf{E}_{X_{n-1}}[f_i(X_n)\mathbf{I}_{D_n} + v_i(X_{t(n)})\mathbf{I}_{\overline{D_n}}]\mathbf{I}_{\{t(\sigma_{\{n\}}) \geq n\}}\} \\
&= \mathbf{E}_x v_i(X_{n-1})\mathbf{I}_{\{t(n-1) \geq n\}}.
\end{aligned}$$

Hence

$$\begin{aligned}
\mathbf{E}_x f_i(X_{t(n)}) &\leq \mathbf{E}_x[f_i(X_{t(n-1)})\mathbf{I}_{\{t(n-1) < n\}} + v_i(X_{n-1})\mathbf{I}_{\{t(n) \geq n\}}] \\
&= \mathbf{E}_x f_i(X_{t(n-1)}).
\end{aligned}$$

We have

$$\mathbf{E}_x f_i(X_{t(n)}) \leq \mathbf{E}_x f_i(X_{t(n-1)}) \leq \mathbf{E}_x f_i(X_{t(0)}) = \mathbf{E}_x f_i(X_{t^*}).$$

We show that if player $j = 1, 2, \dots, p$, $j \neq i$, uses the strategies $^*\sigma^j$ and player i any strategy σ^i then $\mathbf{E}_x f_i(X_{t(\tilde{\sigma})}) \leq \mathbf{E}_x f_i(X_{t(^*\sigma)})$, where

$$\tilde{\sigma} = (^*\sigma^1, \dots, ^*\sigma^{i-1}, \sigma^i, ^*\sigma^{i+1}, \dots, ^*\sigma^p)^T.$$

Let us consider the difference

$$\begin{aligned}
& \mathbf{E}_x f_i(X_{t(\tilde{\sigma})}) - \mathbf{E}_x f_i(X_{t(n)}) \\
&= \mathbf{E}_x [f_i(X_{t(\tilde{\sigma})}) - f_i(X_{t(n)})] \mathbf{I}_{\{t(n) > n\}} \\
&= \mathbf{P}_x \{t(\tilde{\sigma}) > n\} \left\{ \mathbf{E}_x [(f_i(X_{t(\tilde{\sigma})}) - f_i(X_{t(n)})) \mathbf{I}_{\{t(\tilde{\sigma}) > n\}}] - \mathbf{E}_x [v_i(X_{t(n)})] \mathbf{I}_{\{t(n) > n\}} \right\} \\
&\leq \mathbf{P}_x \{t(\tilde{\sigma}) > n\} \left\{ \mathbf{E}_x [\sup_{k \geq n} (f_i^+(X_k)) - v_i(x)] \right\}
\end{aligned}$$

for every $n \in \mathbf{N}$. Since we consider SS from $\mathcal{S}_f^*$ then

$$\mathbf{E}_x f_i(X_{t(\tilde{\sigma})}) - \mathbf{E}_x f_i(X_{t^*}) \leq 0.$$

□

5 Examples

Consider the two person voting game on observation of a homogeneous Markov chain with the majority rule defined in Example 2.1. We have then $p = 2$ and two cases $r = 1$ and $r = 2$. Assume that the Markov chain $(X_n, \mathcal{F}_n, \mathbf{P}_x)$, $n = 0, 1, 2, \dots$ with the state space $\mathbf{E} = \{0, 1\} \times \{0, 1\}$ and $\mathbf{P}_x = \mathbf{P}_{(x_1, x_2)} = \mathbf{P}_{x_1} \times \mathbf{P}_{x_2}$, where $\mathbf{P}_{x_i}$ is given by the transition probability matrix

$$\begin{bmatrix} p & q \\ q & p \end{bmatrix}$$

$p \geq 0, q \geq 0, p + q = 1$. The payoff functions of the players are $f_i(x) = f_i(x_1, x_2) = x_i, i = 1, 2$. Let us formulate the optimality equations and their solution for both cases.

[r=1] The Nash values of the game $(v_1(x), v_2(x))$ fulfill the following equation

$$\begin{aligned}
v_1(x) &= \mathbf{E}_x (f_1(X_1) - v_1(X_1))^+ \mathbf{I}_{\{f_2(X_1) < v_2(X_1)\}} \\
&\quad + \mathbf{E}_x (f_1(X_1) - v_1(X_1)) \mathbf{I}_{\{f_2(X_1) \geq v_2(X_1)\}} + \mathbf{E}_x v_1(X_1)
\end{aligned}$$

$$\begin{aligned}
v_2(x) &= \mathbf{E}_x (f_2(X_1) - v_2(X_1))^+ \mathbf{I}_{\{f_1(X_1) < v_1(X_1)\}} \\
&\quad + \mathbf{E}_x (f_2(X_1) - v_2(X_1)) \mathbf{I}_{\{f_1(X_1) \geq v_1(X_1)\}} + \mathbf{E}_x v_2(X_1).
\end{aligned}$$

The solution of these equations are the functions

$$v_1(x_1, x_2) = v_2(x_2, x_1) = \begin{cases} q/(1 - p^2) & \text{if } (x_1, x_2) = (0, 0) \\ (q + 2pq)/(1 + p) & \text{if } (x_1, x_2) = (0, 1) \\ (2p)/(1 + p) & \text{if } (x_1, x_2) = (1, 0) \\ (q^2 + p^2 + p)/(1 + p) & \text{if } (x_1, x_2) = (1, 1) \end{cases}$$

and the stopping sets are $C^1 = \{(1, 0), (1, 1)\}$ and $C^2 = \{(0, 1), (1, 1)\}$. It is easy to check that for $p \geq q$ we have $v_1(1, 0) \geq v_1(0, 1) \geq v_1(0, 0) \geq v_1(1, 1)$. For $p = q = 1/2$ the equilibrium values $v_1(x) = v_2(x) = 2/3$ for both players are equal and independent of initial state of the Markov chain.

[r=2] In this case the equilibrium values fulfill equations

$$w_1(x) = \mathbf{E}_x(f_1(X_1) - w_1(X_1))^+ \mathbf{I}_{\{f_2(X_1) \geq w_2(X_1)\}} + \mathbf{E}_x w_1(X_1)$$

$$w_2(x) = \mathbf{E}_x(f_2(X_1) - w_2(X_1))^+ \mathbf{I}_{\{f_1(X_1) \geq w_1(X_1)\}} + \mathbf{E}_x w_2(X_1).$$

The solution of these equations are $w_1(x) = w_2(x) = 1$ for every $x \in \mathbf{E}$ and the stopping sets are $C^1 = C^2 = \{(1, 1)\}$.

Acknowledgments. The authors would like to thank Masami Kurano and Jun-ichi Nakagami for their useful discussion. Also we thank to Minoru Sakaguchi for his encouragement and to Lin C. Thomas for his correcting orthographical errors.

References

- [1] Melvin Dresher. *The Mathematics of Games of Strategy. Theory and Applications*. Dover Publications, Inc., New York, 1981.
- [2] P.C. Fishburn. The theory of representative majority decision. *Econometrica*, 39:273 – 284, 1971.
- [3] J.P. Gilbert and F. Mosteller. Recognizing the maximum of a sequence. *J. Amer. Statist. Assoc.*, 61(313):35–73, 1966.
- [4] J.W. Hołubiec and J.W. Mercik. *Inside voting procedures*. Accedo Verlag Gesellschaft, Monachium, 1994.
- [5] J.B. Kadane. Reversibility of a multilateral sequential game: proof of conjecture of Sakaguchi. *J. Oper. Res. Soc. Jap.*, 21(4):509–516, 1978.
- [6] M. Kurano, M. Yasuda, and J. Nakagami. Multi-variate stopping problem with a majority rule. *J. Oper. Res. Soc. Jap.*, 23:205–223, 1980.
- [7] H. Moulin. *Game Theory for the Social Sciences*. New York University Press, New York, 1986.
- [8] J. Nash. Non-cooperative game. *Annals of Mathematics*, 54(2):286–295, 1951.
- [9] Eh.L. Presman and I.M. Sonin. Equilibrium points in game related to the best choice problem. *Theory of Probab. and its Appl.*, 20:770–781, 1975.

- [10] M. Sakaguchi. Optimal stopping in sampling from a bivariate distribution. *J. Oper. Res. Soc. Jap.*, 16(3):186–200, 1973.
- [11] M. Sakaguchi. A bilateral sequential game for sums of bivariate random variables. *J. Oper. Res. Soc. Jap.*, 21(4):486–507, 1978.
- [12] A.N. Shiryaev. *Optimal Stopping Rules*. Springer-Verlag, New York, Heidelberg, Berlin, 1978.
- [13] K. Szajowski. Gra o sumie niezerowej z zatrzymywaniem procesu Markowa. Preprint I18/P-210/84, 1984.
- [14] K. Szajowski. Wielowymiarowe zatrzymywanie procesu Markowa. Preprint I18/P-209/84, 1984.
- [15] A.D. Taylor. *Mathematics and Politics*. Springer, Heidelberg, 1995.
- [16] M. Yasuda, J. Nakagami, and M. Kurano. Multi-variate stopping problem with a monotone rule. *J. Oper. Res. Soc. Jap.*, 25:334–350, 1982.