



# A stochastic game approach for competition over popularity in social networks

Eitan Altman

## ► To cite this version:

Eitan Altman. A stochastic game approach for competition over popularity in social networks. Dynamic Games and Applications, 2013, Special issue on Stochastic Games, 3 (2), pp.313-323. 10.1007/s13235-012-0057-4 . hal-00681959

**HAL Id: hal-00681959**

**<https://inria.hal.science/hal-00681959>**

Submitted on 23 Mar 2012

**HAL** is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

# A stochastic game approach for competition over popularity in social networks

Eitan Altman  
INRIA Sophia-Antipolis,  
2004 Route des Lucioles,  
06902 Sophia-Antipolis Cedex, France  
email: Eitan.Altman@inria.fr

March 21, 2012

## Abstract

The global Internet has enabled a massive access of internauts to content. At the same time it allowed individuals to use the Internet in order to distribute content. When individuals pass through a content provider to distribute contents, they can benefit from many tools that the content provider has in order to accelerate the dissemination of the content. These include caching as well as recommendation systems. The content provider gives preferential treatment to individuals who pay for advertisement. In this paper we study competition between several contents, each characterized by some given potential popularity. We answer the question of when is it worthwhile to invest in advertisement as a function of the potential popularity of a content as well as its competing contents, who are faced with a similar question. We formulate the problem as a stochastic game with a finite state and action space and obtain the structure of the equilibria policy under a linear structure of the dissemination utility as well as on the advertisement costs. We then consider open loop control (no state information) and solve the game using a transformation into a differential game with a compact state space.

## 1 Introduction

We consider in this paper competition between individuals who create contents and wish to propagate the content using some content provider. We assume that an individual can pay the content provider to receive a preferential treatment to his content and have its rate of propagation increased.

As an example, observe Fig ?? that shows the computer screen that I had when watching a video clip on music by Piazzola using Youtube. One can observe three types of advertisements. There is an advertisement for EFS at the bottom of the large dark rectangle which is the screen that shows the video. If one wishes to watch the video then the dark rectangle will occupy the whole computer screen and then this advertisement will be the only one you would see. There is a second advertisement at the top right part of the screen - for courses in Piano Jazz. The first two advertisements just mentioned are not advertisements for content (but they consist a sufficiently important income for youtube so that it can make profits from the free service of displaying video clips). Then to the right we see the first five video clips in a recommendation list provided by google. The first in the list has a tag "Ad". It is a

video clip that received a priority in the recommendation list. The remaining clips in the recommendation list did not have to pay anything.

When some content makes it to the first ones in the list then it gets a higher visibility than the others and therefore the speed of propagation is expected to increase.

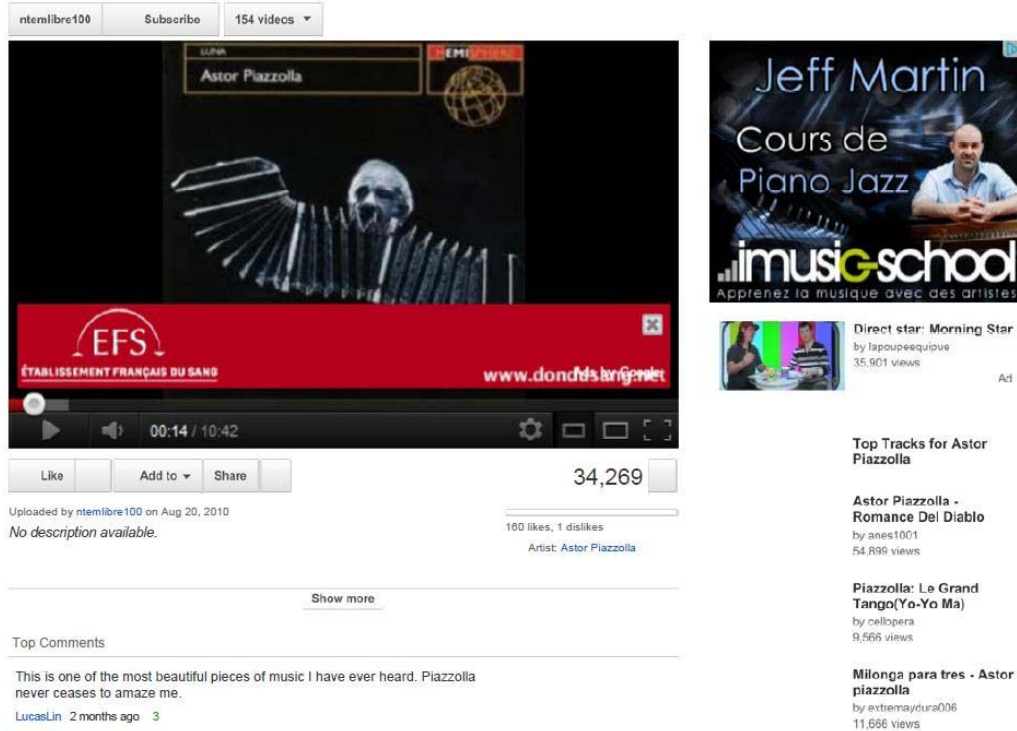


Figure 1: Publicity in Youtube

We consider a competition between several contents, each having possibly another level of potential popularity (or in other words, another rate of propagation.) Depending on the popularity level of the contents, on the potential size of the interested audience as well as the number of past downloads of each of the contents, each individual may decide whether or not to purchase a higher priority. We formulate this decision problem as a stochastic game with a finite state and action spaces. The solution of the problem allows us to provide guidelines for individual's advertisement strategies.

We formulate the problem as a continuous time Markov game. We then use uniformization in order to transform the problem into an equivalent discrete time Markov game. In the case of linear costs we manage to reduce considerably the dimension of the state space and obtain a characterization of the equilibrium policy.

The structure of the paper is the following. The next section provides the problem statement and the stochastic game model. It is introduced as a continuous time Markov game. We transform it in Section 3 into a discrete time finite state and action stochastic game. We obtain the structure of the equilibrium in Section 4. In section 5 we transform the problem

into a deterministic equivalent dynamic game and show how to reduce its dimensionality. We then solve in Section 6 the problem with infinite horizon criterion. We end with a short concluding section.

## 2 Model and statement of the problem

Assume that there are  $N$  competing contents (say some softwares that are sold over the Internet). There are  $M$  potential common destinations. We assume that a destination wishes to acquire one of these contents and will purchase the one at the first possible opportunity.

We assume that opportunities for purchasing a content  $n$  arrive at destination  $m$  according to a Poisson process with parameter  $\lambda_n$  starting at time  $t = 0$ . Hence if at time  $t = 0$  destination  $m$  wishes to purchase the content  $n$ , it will have to wait some time which is exponentially distributed with parameter some parameter  $\lambda_i$ .

The value of  $\lambda_i$  may differ from one content to another. The difference is partly due to the fact that different contents may have different popularity.

We assume that the owner of a content  $n$  can accelerate the propagation speed of the propagation of the content in two ways: First, it can increase  $\lambda_i$  by some advertisement effort.

Secondly, we allow for content  $i$  to be available for a subset of  $x_i(0)$  destination at time 0 (without waiting for a purchase opportunity). This again can be achieved using some advertisement effort. Here are some examples. When selling books, it is often possible to command a book even before it appears. Advertisements of movies, concerts, theatre and other cultural events, as well as sport events often begins quite before the opening and one can then purchase tickets way before the premier.

We next model the problem as a continuous time Markov game.

- **State space.** Let  $x_i(t)$  be the number of destinations that have content  $i$  at time  $t$ . It is a continuous time Markov chain with a finite state space  $\mathbf{X} = \{(x_1, \dots, x_N) \in N^m, \sum_{i=1}^N x_i \leq M\}$ .
- **Action Space.** Let  $\mathbf{A}_i$  be a finite set of actions available to the owner of content type  $i$ .  $a \in \mathbf{A}_i$  is a possible value of the amount of acceleration of  $\lambda_i$ . Let  $\mathbf{A}$  be the product action space of  $\mathbf{A}_i$ ,  $i = 1, \dots, N$ .  
For all  $i$ , any action  $a \in \mathbf{A}_i$  satisfies  $a \geq 1$ . The action  $\underline{a} = 1$  is the one that does not use any acceleration. Let  $a_1$  be the smallest action not including  $\underline{a}$  and let  $\bar{a}$  denote the largest action.
- **The transition intensity.** Let  $|\mathbf{n}| := \sum_{i=1}^N n_i$ . Given that the state at that time is  $\mathbf{n}$  and the action of players is  $\mathbf{a} \in \mathbf{A}$ , the transition intensity is given by

$$Q(\mathbf{n} + \mathbf{e}_i | \mathbf{n}) = \lambda_i a_i (M - |\mathbf{n}|).$$

where  $\mathbf{e}_i$  is the unit vector whose  $i$ th component equals 1 and the rest are zero. Indeed, at state  $\mathbf{n}$ , the number of destinations that do not yet have any content is given by  $M - |\mathbf{n}|$ . the time till the first one of these receives the content of type  $i$  is the minimum of  $M - |\mathbf{n}|$  independent exponential random variables each with parameter  $\lambda_i a_i$ . It is thus an exponential random variable with parameter  $\lambda_i a_i (M - |\mathbf{n}|)$ .

- **Policies.** A pure stationary policy for player  $i$  is a map from  $\mathbf{X}$  to  $\mathbf{A}_i$ . Let  $\Delta(\mathbf{A}_i)$  be the set of probability measures over  $\mathbf{A}_i$ . A mixed stationary policy is a map from  $\mathbf{X}$  to  $\Delta(\mathbf{A}_i)$ . Choose some horizon  $T$ . A Markov policy for player  $i$  is a measurable function  $w^i$  that assigns for each  $t \in [0, T]$  and each state  $\mathbf{x}$  a mixed action  $w_t^i(\mathbf{x})$ . For a given initial state  $\mathbf{x}$  and a given Markov policy  $w$ , there exists a unique probability measure  $P_{\mathbf{x}}^w$  which defines the state and action random processes  $X(t), A(t)$ . Multi-policies are defined as vectors of policies, one for each player.
- **The utility.** Let  $c_i(a_i) = \gamma_i(a_i - 1)$  be the cost for player  $i$  of choosing action  $a_i$ . Since  $a_i$  has the interpretation of the factor by which the player wishes to accelerate the dissemination,  $c_i(a_i)$  is increasing in its argument and when there is no acceleration ( $a_i = \underline{a} = 1$ ) the cost is zero. The utility is assumed to be a weighted sum of a payoff that is proportional to the expected number of destinations that have the content at time  $T$ , and some disutility that describes the total advertisement cost:

$$U_i(T) = E[X_i(T)] - E \left[ \int_0^T c_i(A_i(t)) dt \right] = E \left[ \int_0^T -c_i(A_i(t)) dt + dX_i(t) \right]$$

(Note that  $X_i(t)$  is monotone so that the integral is well defined).

We shall obtain interesting structure of the optimal policy. To understand the reasons for that, we shall first present a more general utility function. We replace the instantaneous cost  $\gamma_i a_i$  by a general one of the form  $c_i(A_i(t))$ , and we replace the final expected dissemination utility  $E[X_i(T)]$  with  $E[g_i(X_i(T))]$ .

Define  $\zeta_i(m) = g_i(m+1) - g_i(m)$ .

Let  $\mathcal{M}$  denote the set of states at which  $\sum_{i=1}^M x_i = M$ , i.e. all states at which all destinations have purchased the content. Every state in the set  $\mathcal{M}$  is an absorbing state. The stochastic game is absorbing, and under any policy  $w$ , the time to absorption is finite  $P_w$  a.s. It is moreover, stochastically smaller than the one under the policy  $\mathbf{0}$  in which no player ever accelerates. All states other than  $\mathcal{M}$  are transient. Once  $\mathcal{M}$  is reached, no player has any incentive to ever accelerate; we may assume without loss of generality that only  $\underline{a}$  is available for states in  $\mathcal{M}$ . The utility for each player  $i$  can then also be written as

$$U_i(T) = E \left[ \int_0^{\min(T, \sigma)} -c_i(A_i(t)) dt + dX_i(t) \right]$$

where  $\sigma$  is the hitting time of the set  $\mathcal{M}$ . The Markov game has thus a structure of an absorbing Markov Decision Process (MDP), see [1, chap 7].

When considering the game within a finite horizon then we shall restrict our search of equilibrium to the Markovian multi-policies. For the infinite horizon problem, we shall restrict to stationary mixed multi-policies and find equilibria within this class.

Indeed, for a finite horizon, if we obtain an equilibrium within Markov policies then at equilibrium, each player is faced with an absorbing MDP for which there exists an optimal Markov policy. Thus no player can benefit by using any other more general policy (see [1]). A similar argument shows that for infinite horizon, we can restrict to stationary policies.

### 3 Uniformization

We consider the game in which the past and present state as well as the past actions are known to all players. We use the standard uniformization approach [3] to transform the decision problem into a discrete time Markov game.

Introduce the following discrete time Markov Game. The state and action spaces are the same. The transition probabilities are defined as follows. Define  $\lambda = M \sum_i \lambda_i \bar{a}_i$ .

$$P_{\mathbf{x} \mathbf{a} \mathbf{z}} = \begin{cases} (M - |\mathbf{x}|) \frac{a_i \lambda_i}{\lambda} & \text{for } \mathbf{z} = \mathbf{x} + e_i, \mathbf{x} \in \mathbf{X} \setminus \mathcal{M} \\ 1 - (M - |\mathbf{x}|) \frac{\sum_i a_i \lambda_i}{\lambda} & \text{for } \mathbf{z} = \mathbf{x}, \mathbf{x} \in \mathbf{X} \end{cases} \quad (1)$$

Define

$$\delta_j(v, \mathbf{x}) = v(\mathbf{x} + e_j) - v(\mathbf{x}).$$

Define for each player  $i$ ,  $\mathbf{x} \in \mathbf{X} \setminus \mathcal{M}$ ,  $\mathbf{a} \in \mathbf{A}$  and  $v \in R^{\mathbf{X}}$ :

$$J^i(v, \mathbf{x}, \mathbf{a}) = -c_i(a_i) + \frac{(M - |\mathbf{x}|)}{\lambda} \sum_{j=1}^N a_j \lambda_j (\zeta_i(x_i) 1\{j = i\} + \delta_j(v, \mathbf{x}))$$

and set  $J^i(v, \mathbf{x}, \mathbf{a}) = 0$  for  $x \in \mathcal{M}$ .  $J^i(v, \mathbf{x}, \mathbf{a})$  is the total utility for player  $i$  if at time 0 the system is at state  $\mathbf{x}$ , player  $j$  takes action  $a_j$  (where  $a_j$  is the  $j$ th component of the action vector  $\mathbf{a}$ ) and the utility to go for player  $i$  from the next transition onwards is  $v(\mathbf{y})$  if the state after the next transition is  $\mathbf{y}$ .

Let  $\mathbf{u}$  be a mixed stationary multi-policy. With some abuse of notation we define for each player  $i$  and for each  $\mathbf{x} \in \mathbf{X} \setminus \mathcal{M}$ ,

$$J^i(v, \mathbf{x}, \mathbf{u}) = - \sum_{a \in \mathbf{A}_i} u_i(a|\mathbf{x}) c_i(a) + \frac{(M - |\mathbf{x}|)}{\lambda} \sum_{j=1}^N \left[ \sum_{a \in \mathbf{A}_j} u_j(a|\mathbf{x}) a \right] \lambda_j (\zeta_i(X_i(t)) 1\{j = i\} + \delta_j(v, \mathbf{x}))$$

and set  $J^i(v, \mathbf{x}, \mathbf{u}) = 0$  for  $x \in \mathcal{M}$ .

$$0 = \max_{u \in \Delta(\mathbf{A}_i)} J^i(v_i, \mathbf{x}, \mathbf{u}) \quad (2)$$

**Theorem 1.** (i) The fixed point equation (2) has a solution  $v^*$ .

(ii) Let  $v^*$  be such a fixed point. Any mixed stationary multi-policy  $\mathbf{u}$  such that achieves the argmax of (2) for all  $i$  is a mixed stationary Nash equilibrium.

**Proof.** A similar proof is already available for the discounted cost criterion, and under some additional assumptions, in the case of the average reward problem (see e.g. [2]). The proof in our case follows the same steps. The only step that is not direct is the continuity of the performance measures in the stationary policies.

We first note that this Markov game is absorbing: it has an absorbing set that is reached under any policy with probability 1 and the expected time to hit the set is uniformly bounded over all policies. (For more details, see discussion in the Concluding Section.) The required continuity then follows known results (see e.g. [1]). ■

## 4 The game on an aggregated state space

In the next subsection we shall go back to the original linear structure of the utilities. We begin by considering that dissemination utility is linear with the form  $g_i(x_i) = x_i$ .

We then consider that also the acceleration costs  $c_i$  are linear and have the form  $c_i(a_i) = \gamma_i(a_i - 1)$  for some constants  $\gamma_i > 0$ .

### 4.1 Linear dissemination utilities

We present a surprisingly simple structure of the equilibrium policy for the case of linear dissemination utility. We show that one can transform the stochastic game into an equivalent one which has the same action space but a much simpler state space: it is one dimensional and is given by the set  $\bar{\mathbf{X}} = \{0, 1, \dots, M\}$ .

Define  $\bar{\mathbf{X}} = \{0, 1, \dots, M\}$  to be the class of aggregated states. An aggregated state  $i \in \bar{\mathbf{X}}$  corresponds to the set of states  $\mathbf{x} \in \mathbf{X}$  such that  $|\mathbf{x}| = i$ . An aggregated state thus counts the total number of destinations that have some content. Taking the summation in (1) we get the following transition probabilities for the aggregated Markov game:

$$P_{xaz} = \begin{cases} (M - x) \frac{\sum_{i=1}^N a_i \lambda_i}{\sum_{i=1}^N a_i \lambda_i} & \text{for } z = x + 1, x \in \bar{\mathbf{X}} \setminus \{M\} \\ 1 - (M - x) \frac{\sum_{i=1}^N a_i \lambda_i}{\sum_{i=1}^N a_i \lambda_i} & \text{for } z = x, x \in \bar{\mathbf{X}} \end{cases} \quad (3)$$

The aggregated state process has the Markov property: the dependence of the next aggregated state on the history is only through the current aggregated state and actions. However, the dissemination instantaneous utility,  $\zeta_i(x_i)$  cannot be written as a function of the aggregated utility.

We shall consider in this section the original dissemination utility  $g_i(x_i) = x_i$ . We thus get  $\zeta_i(x_i) = 1$ . Hence when using the equivalent instantaneous dissemination utility, it is no more a function of the state. We thus get a Markov game formulation with a considerably reduced complexity. Any equilibrium in this new stochastic game is also an equilibrium in the original one. (Indeed, this follows from Theorem 6.3 in [1]).

### 4.2 Computing the equilibrium

Now that we reduced the state space to  $M + 1$  states only (of which state  $M$  is absorbing) it remains to compute for each of these states the randomized action of each user at equilibrium

Fix some stationary policy  $u$ . Let  $X(t) = \sum_{i=1}^N X_i(t)$ . Define for  $m = 0, \dots, M - 1$  the total expected reward from the moment that  $X(t) = m$  till it reaches  $m + 1$  by  $U_i^m(\mathbf{u})$ . We note that the time until  $X(t)$  jumps from  $m$  to  $m + 1$  is an exponentially distributed random variable with parameter

$$\theta_m(\mathbf{a}) = (M - m) \sum_{j=1}^N a_j \lambda_j$$

The probability that the transition to  $j + 1$  occurred due to player  $i$  is given by

$$p_i = \frac{a_i \lambda_i}{\sum_{j=1}^N a_j \lambda_j}$$

Hence

$$U_i^m(\mathbf{a}) = \frac{c_i(a_i)}{\theta_m} + p_i(\mathbf{a}) = \frac{c_i(a_i) + (M - m)a_i\lambda_i}{(M - m) \sum_{j=1}^N a_j\lambda_j}$$

We conclude that a stationary mixed equilibrium can be obtained as follows:

**Theorem 2.** *Consider the case of linear dissemination utility. Denote by  $\mathbf{u}^*(m)$  an equilibrium multi-strategy in the  $m$ th matrix game,  $m = 0, \dots, M - 1$ , in which the utility of player  $i$  is given by  $U_i^m(\mathbf{a})$ . Then the mixed stationary policy for which each player  $i$  chooses an action  $a$  with probability  $u^*(a|m)$  whenever the state satisfies  $|\mathbf{m}| = m$ , is an equilibrium for the original problem.*

### 4.3 Linear acceleration costs

Assume next that for some  $i$ ,  $c_i(a_i) = \gamma_i(a_i - 1)$  for some constants  $\gamma_i$ . Define  $\Delta_i(m) = -\gamma_i + (M - m)\lambda_i$ . Then

$$\begin{aligned} U_i^m(\mathbf{a}) &= \frac{a_i(-\gamma_i + (M - m)) + \gamma_i}{(M - m) \sum_{j=1}^N a_j\lambda_j} \\ &= \frac{1}{\lambda_i(M - m)} \frac{(-\gamma_i + (M - m)\lambda_i) a_i\lambda_i + \gamma_i\lambda_i}{\sum_{j=1}^N \lambda_j a_j} \\ &= \frac{1}{\lambda_i(M - m)} \left( -\gamma_i + (M - m)\lambda_i - \frac{\Delta_i^m}{\sum_{j=1}^N \lambda_j a_j} \right) \end{aligned}$$

where

$$\Delta_i^m = (-\gamma_i + (M - m)\lambda_i) \sum_{j \neq i} \lambda_j a_j - \gamma_i\lambda_i$$

Then for any action of players  $j \neq i$ , the following holds. if  $\Delta_i(m) > 0$  then  $U_i^m(\mathbf{a})$  is maximized at  $\bar{a}_i$ . Otherwise it is maximized at  $\underline{a} = 1$ .

Since  $\Delta_i(m)$  is increasing in  $m$ , then if  $\Delta_i(m) > 0$  for some  $m$  then  $\Delta_i(j) > 0$  for all  $j > m$ . Thus if  $\Delta_i(m) > 0$  then for all  $j \geq m$ , the utility of player  $i$  is maximized at  $a_i = \bar{a}_i$ .

We conclude that if for some  $i$ ,  $c_i(a_i) = \gamma_i a_i$ , then at equilibrium, player  $i$  has a threshold policy  $L_i$ : it uses  $\bar{a}$  at all states above  $L_i$  and  $\underline{a} = 1$  otherwise.  $L_i$  is given by the smallest integer greater than or equal to  $\rho_i$ , where  $\rho_i$  is the solution of  $0 = -\gamma_i + (M - m)\lambda_i$  and is thus given by

$$\rho_i = M - \frac{\gamma_i}{\lambda_i} \quad (4)$$

In particular, if  $\rho_i \leq 0$  then at all states the equilibrium policy uses the largest acceleration available,  $\bar{a}_i$ , and if  $\rho_i > M$  then the equilibrium policy for player  $i$  always uses no acceleration.

## 5 The case of no state information

We shall assume below that the players

- do not observe the state.
- either know the initial state or know its distribution or its expectation. All players are assumed to have the same information on the expected value of  $X_0$ .

We shall therefore restrict to the subset of Markov policies that depend on time and on the available information on the initial state, but not on the state at time  $t > 0$ . We denote the set of such policies for player  $i$  by  $\mathbf{W}^i$ .

We show next that the stochastic game is equivalent to a differential game.

We assume that all players know the policies used by other players as well as the value of  $\bar{x}_i(0)$ ,  $i = 1, \dots, N$ .

Fix a policy  $w \in \mathbf{W}$ . Let  $\bar{x}_i(t) := E[X_i(t)]$  and  $\bar{x}(t) := \sum_{i=1}^N \bar{x}_i(t)$ . Then

$$\dot{\bar{x}}_i(t) = \lambda_i w_t^i (M - \bar{x}(t)) \quad (5)$$

The utility of player  $i$  is given by  $U_i(T, w, z)$  where

$$U_i(t, w, z) = \bar{x}_i(t) - \int_0^t c_i(w_s^i) ds$$

where  $\bar{x}(0) = z$ . We thus obtained a differential game. Note that although  $w \in \mathbf{W}$  does not have knowledge of the realization of the state trajectory, we can allow  $w_t^i$  to depend on  $\bar{x}(t)$  since each player can compute it from the knowledge of the policies used and from the knowledge of the expected initial states.

This is an  $N$ -dimensional differential game. We shall next transform it into an equivalent one-dimensional problem where  $y(t) = M - \bar{x}(t)$  is the state.

Indeed, we show that both the dynamics as well as the utilities can be written directly in terms of the state trajectory  $y(t)$ . Taking the summation over  $i$  in (5), we get,

$$\dot{y}(t) = -\dot{\bar{x}}(t) = -\eta_t y(t) \quad (6)$$

where  $\eta_t = \sum_{i=1}^N \lambda_i w_t^i$ ,  $y(0) = M - \sum_{i=1}^N \bar{x}_i(0)$ .

Moreover, the utility can be written as  $U_i(T, w, z)$  where

$$U_i(t, w, z) = \bar{x}_i(0) + \int_0^t [-c(w_s^i) + \dot{\bar{x}}_i(s)] ds = \bar{x}_i(0) + \int_0^t r(w_s, y_s) ds$$

where

$$r(a_i, y) = -c(a_i) + \lambda_i a_i y.$$

It is indeed a function of the trajectories of  $y_t$  and  $a_t$  only; note that  $\bar{x}_i(0)$  are constants that are not affected by the decisions of the players. The last equality was obtained by substituting (6).

**Remark 1.** (i) The solution of (6) is

$$y(t) = (M - x(0)) \left( 1 - \exp \left( - \int_0^t \eta(s) ds \right) \right)$$

Thus  $x_i(t)$  is the solution of

$$\dot{x}_i = w_t^i \lambda_i y(t) \quad (7)$$

(ii) Under any policy,

$$y(t) \leq (M - x(0)) \exp \left( -t \sum_{i=1}^N \lambda_i \right). \quad (8)$$

## 6 Infinite horizon with no information

We consider the total cost problem is (i.e. the game obtained for an infinite horizon). At equilibrium, the utility for each player  $i$  should be the value of the best response policy against the others' policies. The value for player  $i$  is known to be the unique viscosity solution of the following provided that it is piecewise differentiable.

$$0 = \sup_{a_i \in \Delta(\mathbf{A}^i)} J^i(\mathbf{a}, y)$$

where

$$J^i(\mathbf{a}, y) = \left( r_i(a_i, y) - \dot{v}_i(y) y \sum_{j=1}^N \lambda_j a_j \right)$$

We shall compute explicitly the equilibrium below.

Assume that on some neighbourhood of  $y$ , the sup is achieved by some vector  $b$ . Then

$$\begin{aligned} \dot{v}_i(y) &= \frac{r_i(b_i, y)}{y \sum_{j=1}^N \lambda_j b_j} = \frac{-c(b_i) + \lambda_i b_i y}{y \sum_{j=1}^N \lambda_j b_j} \\ &= \frac{-c(b_i)}{y \sum_{j=1}^N \lambda_j b_j} + \frac{\lambda_i b_i}{\sum_{j=1}^N \lambda_j b_j} \end{aligned}$$

Thus

$$v_i(y) = v_i(s) - (\log(y) - \log(s)) \frac{c(b_i)}{\sum_{j=1}^N \lambda_j b_j} + \frac{\lambda_i b_i (y - s)}{\sum_{j=1}^N \lambda_j b_j}$$

We also have

$$J^i([a_i, \mathbf{b}^{-i}], y) = r_i(a_i, y) - r_i(b_i, y) \frac{\sum_{j \neq i} \lambda_j b_j + \lambda_i a_i}{\sum_{j=1}^N \lambda_j b_j}$$

An action maximises this expression (over  $a_i$ ) if and only if it maximizes

$$\frac{r_i(a_i, y)}{\sum_{j \neq i} \lambda_j b_j + \lambda_i a_i} - \frac{r_i(b_i, y)}{\sum_{j=1}^N \lambda_j b_j}$$

which is equivalent to maximizing

$$V([a_i, \mathbf{b}^{-i}], y) := \frac{r_i(a_i, y)}{\sum_{j \neq i} \lambda_j b_j + \lambda_i a_i}$$

This implies the following.

**Theorem 3.** Consider  $N$ -player matrix games where the utility of player  $i$  is given by

$$\bar{V}^i(\mathbf{a}, y) = \frac{r_i(a_i, y)}{\sum_{j=1}^N \lambda_j a_j}.$$

$y$  is a parameter taking values in  $[0, M]$ . Let  $\mathbf{u}(y)$  denote a mixed equilibrium in the matrix game  $y$ . Then  $\mathbf{u}$  is a stationary equilibrium in the original game.

We obtained the same form of optimal policy as in the discrete case with state information, except that the parameter  $y$  is now an interval rather than the finite set  $\{0, \dots, M-1\}$ .

Assume that  $c_i(a_i) = \gamma_i(a_i - 1)$  for some positive constant  $\gamma_i$ . Then the maximizer  $a_i$  is independent of  $\mathbf{b}^{-i}$ . It is given by  $\underline{a}$  for  $y < \gamma_i/(\lambda_i)$ , and by the maximal element of  $\mathbf{A}_i$  if the converse inequality holds.

We conclude the following.

- At equilibrium, each player  $i$  accelerates the  $\lambda_i$  by the largest possible  $a_i$  as long as  $\bar{x} < \rho_i$  and does not accelerate for  $\bar{x} > \rho_i$ .  $\rho_i$  is the same as the one derived in the discrete case, see (4).
- Assume that  $\gamma_i$  is the same for all  $i$ . Then the owner of a more popular content (i.e. with a larger  $\lambda_i$ ) will advertize over a larger set (interval) of states

## 7 Concluding comments

We comment on the relation between the differential game and the original stochastic game. Every set in  $\mathcal{M}$  is also absorbing in the differential game. However, it is never reached from any other state. However, starting at any state not in  $\mathcal{M}$ , the distance to  $\mathcal{M}$  converges to zero exponentially fast (and uniformly) as it follows from (8).

Although the differential game that we solved is different than the original discrete one (they differ in the information available), it was seen to have a similar structure of equilibria in the case of linear dissemination utility. One can show in fact that the differential game is a fluid limit for the discrete game with a proper scaling of the state space and of the rates  $\lambda_i$ 's.

This work is a first step for us in understanding competition issues between content producers over the Internet. Many other aspects will be modeled in the future, including coupling between various social networks; indeed, one way of accelerating the dissemination in the Internet of, say, some movie, would be to advertize it using Twitter and Facebook. Coupling can occur by using the "sharing" option which allows to migrate a content from one social network to another.

## References

- [1] E. Altman. *Constrained Markov Decision Processes*. Chapman and Hall/CRC, 1999.
- [2] A. Federgruen. On n-person stochastic games with denumerable state space. *Adv. Appl Prob.*, 10:452–471, 1978.
- [3] S.A. Lippman. Applying a new device in the optimization of exponential queueing systems. *Operations Research*, 23:687–710, 1975.