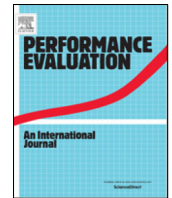




Contents lists available at ScienceDirect

Performance Evaluation

journal homepage: www.elsevier.com/locate/peva

When to arrive at a queue with earliness, tardiness and waiting costs

Eliran Sherzer*, Yoav Kerner

Industrial Engineering, Ben-Gurion University, Beer-Sheva, Israel

ARTICLE INFO

Article history:

Received 20 September 2016

Received in revised form 27 August 2017

Accepted 29 August 2017

Available online xxxx

ABSTRACT

We consider a queueing facility where customers decide when to arrive. All customers have the same desired arrival time (w.l.o.g. time zero). There is one server, and the service times are independent and exponentially distributed. The total number of customers that demand service is random, and follows the Poisson distribution. Each customer wishes to minimize the sum of three costs: earliness, tardiness and waiting. We assume that all three costs are linear with time and are defined as follows. Earliness is the time between arrival and time zero, if there is any. Tardiness is simply the time of entering service, if it is after time zero. Waiting time is the time from arrival until entering service. We focus on customers' rational behavior, assuming that each customer wants to minimize his total cost, and in particular, we seek a symmetric Nash equilibrium strategy. We show that such a strategy is mixed, unless trivialities occur. We construct a set of equations that its solution provides the symmetric Nash equilibrium. The solution is a continuous distribution on the real line. We also compare the socially optimal solution (that is, the one that minimizes total cost across all customers) to the overall cost resulting from the Nash equilibrium.

© 2017 Elsevier B.V. All rights reserved.

1. Introduction

In many real life situations, customers choose strategically when to arrive to a queue. For example, when commuting to work in the morning, or going to the cafeteria at lunch time, etc. In those kind of situations, customers usually have a preferred arrival time, such as, commuting to work right after dropping the kids at school. For the sake of simplicity, we refer to the preferred arrival time as time 0 (w.l.o.g.). In many cases, customers share the same preferred arrival time. This is very likely to be the case in the morning commute, or when going out for lunch. If all customers wish to arrive at that time, it causes congestions. As a consequence, one would consider altering his arrival time in order to avoid waiting for a long time. In this situation, from an individual point of view, a trade-off between the arrival time preference and the queue length is formed. This trade-off was studied widely in the literature (see [1–4] and [5]). In those studies, the trade-off was between waiting and tardiness costs. However, we suggest an additional cost, that is also reflected in the dilemma of when to arrive, and it is referred as *earliness cost*. This cost is induced by making the effort of arriving earlier than desired. A classic example would be commuting to work at 4 AM just to avoid traffic jams. In this example, clearly, arriving this early minimizes the waiting time in one hand, but on the other, commuting so early is inconvenient. We further discuss the motivation of adding earliness cost in the next section.

In this study, we propose a model where customers choose when to arrive. Each individual possesses three types of costs: Earliness, tardiness and waiting, where waiting refers to the amount of time one waited in the queue. Earliness is related to

* Corresponding author.

E-mail addresses: eliransc@post.bgu.ac.il (E. Sherzer), kerneryo@bgu.ac.il (Y. Kerner).

the amount of time one arrived before time 0, and if he arrived after time 0, then there is no earliness cost. Tardiness is related to the amount of time entering service after time 0. However, if an individual service time began prior to time 0, then there is no tardiness cost. We focus on homogeneous customers with respect to their costs and desired service time. Under this assumption, customers make their arrival decision while knowing that others have the same preferences. In such systems, all customers' decisions interact. That is, an individual's cost is determined not only by his own decision, but also by others', and thus the natural model of customers' arrivals is a non-cooperative game. Motivated by this point of view, we look for a symmetric Nash equilibrium. That is, if a (possibly) mixed arrival strategy is used by all, an individual has no incentive to unilaterally deviate from it. This solution, which determines the arrival process for the system, implies the opening and closing times of the system. We also discuss a system with exogenous opening and closing times, and study the impact of these constraints on customers' behavior.

We found that in the case where there were no opening and closing time constraints, the Nash equilibrium is mixed and unique, and has a positive density $f(t)$ along the interval $[-T_{e1}, T_{e2}]$, with some $T_{e1}, T_{e2} > 0$. $f(t)$ is complete (that is, the CDF has no jumps) and continuous almost everywhere, with a connected support. In the case where the opening and closing times are constrained, the equilibrium may be either pure or mixed, and if the latter is the case, then it can also contain an atom.

In this paper we show how to compute the Nash equilibrium strategy for all of the above-mentioned cases. The equilibrium strategy, in general, is not socially optimal, in the sense that it does not minimize the expected overall costs of all the customers. This is because the decision of when to arrive, taken by an individual, imposes additional losses on the others. The latter are often referred to as externalities. Therefore, finding the socially optimal solution is also of our interest. This allows us to compare the strategy that holds for the Nash equilibrium and the one that minimizes the overall cost. In order to compare the two, we derive their ratio, which is called the *price of anarchy* (PoA). The value of the PoA can help us understand how much customers can minimize their costs by cooperating with each other.

The analysis in the paper is performed for two types of environments. The first is the familiar stochastic environment. In this case we assume that the number of users is random and follows the Poisson distribution¹ and the service times follow the exponentially distribution, independent and assumed to be work conserving. The second environment is fluid-based, and each user is associated with a drop of infinitesimal size. The stochastic model is more accurate but can provide only numerical results (and not analytical). The fluid model on the other hand can provide analytic results. Moreover, PoA can be obtained as well. We make some observations regarding the relations between the two models later on.

The paper is organized as follows. In Section 2 we describe the model motivation and provide a literature review. In Section 3 we present the model under study. In Section 4, we obtain the arrival equilibrium strategy for both the stochastic model and the fluid model. In addition, we compare via the fluid model the social cost of the equilibrium arrival strategy and the socially optimal cost. In Section 5, we obtain the above solution concepts for the constrained model. Finally, in Section 6 we summarize our main results.

2. Motivation and literature review

We next provide a motivation for the model under study. More specifically, we explain why adding earliness cost to our model is important. This is followed by a literature review.

2.1. Motivation

In the previous section, we introduced the concept of earliness cost. Whereas, earliness cost is induced by making the effort of arriving earlier than desired. We next specify why we think it is vital to include it in the model. For this purpose, we consider the morning commute example. When commuting, it is very natural having a preferred arrival time. For the sake of demonstration, we consider an employee, that does not want to wake up too early for work in one hand, but on the other hand, does not want to arrive too late, when there are important meetings to attend. For this particular example, we set the desired arrival time to be 08:00 AM. The problem is, arriving at 08:00 comes with heavy traffic. As a consequence, each arrival point, is costly, due to the different cost types. In fact, each point comes with an aggregation of the three cost types. For illustration, we consider a few arrival epochs. Arriving between 05:30 and 06:30 AM induces an expected high earliness cost, low waiting cost and no tardiness cost at all, since heavy traffic is not expected and there is no reason to be late in this case. Arriving between 07:30 and 08:30 AM induces a high waiting cost due to heavy traffic, and relatively low earliness and tardiness costs. Whereby, arriving between 10:30 and 11:30 AM, induces an expected high tardiness cost, no earliness cost since no effort to arrive earlier was made, and probably low expected waiting cost as the traffic is probably flowing normally at that point.

The main idea is that when one needs to decide when to arrive, there are three distinctive types of costs to consider. Moreover, they all should be under consideration simultaneously, as they are all time-dependent. To the best of our knowledge, we are the first to identify this additional cost (i.e. earliness) under a game theoretic environment.² We further

¹ One can think of a huge amount of potential users, each one participates with a significantly small probability. The Poisson approximation of the binomial distribution leads to our assumption (see more in [2]).

² Vickrey, already in 1969, considered the earliness concept in [6], but not as a game.

Download English Version:

<https://daneshyari.com/en/article/6888526>

Download Persian Version:

<https://daneshyari.com/article/6888526>

[Daneshyari.com](https://daneshyari.com)