

Polling Models with Two-Stage Gated Service: Fairness Versus Efficiency*

R.D. van der Mei^{1,2} and J.A.C. Resing³

¹ CWI, Advanced Communication Networks
P.O. Box 94079, 1090 GB Amsterdam, The Netherlands

² Vrije Universiteit, Faculty of Sciences
De Boelelaan 1081a, 1081 HV Amsterdam, The Netherlands

³ Eindhoven University of Technology, Mathematics and Computer Science
P.O. Box 513, 5600 MB Eindhoven, The Netherlands

Abstract. We consider an asymmetric cyclic polling system with general service-time and switch-over time distributions with so-called two-stage gated service at each queue, an interleaving scheme that aims to enforce fairness among the different customer classes. For this model, we (1) obtain a pseudo-conservation law, (2) describe how the mean delay at each of the queues can be obtained recursively via the so-called Descendant Set Approach, and (3) present a closed-form expression for the expected delay at each of the queues when the load tends to unity (under proper heavy-traffic scalings), which is the main result of this paper. The results are strikingly simple and provide new insights into the behavior of two-stage polling systems, including several insensitivity properties of the asymptotic expected delay with respect to the system parameters. Moreover, the results provide insight in the delay-performance of two-stage gated polling compared to the classical one-stage gated service policies. The results show that the two-stage gated service policy indeed leads to better fairness compared to one-stage gated service, at the expense of a decrease in efficiency. Finally, the results also suggest simple and fast approximations for the expected delay in stable polling systems. Numerical experiments demonstrate that the approximations are highly accurate for moderately and heavily loaded systems.

1 Introduction

This paper is motivated by dynamic bandwidth allocation schemes in an Ethernet Passive Optical Network (EPON), where packets from different Optical Network Units (ONUs) share channel capacity in the upstream direction. An EPON is a point-to-multipoint network in the downstream direction and a multi-point to point network in the upstream direction. The Optical Line Terminal (OLT) resides in the local office, connecting the access network to the Internet. The OLT allocates the bandwidth to the Optical Network Units (ONUs) located at the

* A preliminary version of this paper has also been presented at the Second Korea-Netherlands Conference on Queueing Theory and its Applications to Telecommunication Systems (Amsterdam, October 24-27, 2006).

customer premises, providing interfaces between the OLT and end-user network to send voice, video and data traffic. In an EPON the process of transmitting data downstream from the OLT to the ONUs is broadcast in variable-length packets according to the 802.3 protocol [6]. However, in the upstream direction the ONUs share capacity, and various polling-based bandwidth allocation schemes can be implemented. Simple time-division multiplexing access (TDMA) schemes based on fixed time-slot assignment suffer from the lack of statistical multiplexing, making inefficient use of the available bandwidth, which raises the need for dynamic bandwidth allocation (DBA) schemes. A dynamic scheme that reduces the time-slot size when there is no data to transmit would allow excess bandwidth to be used by other ONUs. However, the main obstacle of implementing such a scheme is the fact the OLT does not know in advance how much data each ONU has to transmit. To overcome this problem, Kramer et al. [7,8] propose an OLT-based interleaved polling scheme similar to hub-polling to support dynamic bandwidth allocation. To avoid monopolization of bandwidth usage of ONUs with high data volumes they propose an interleaved DBA scheme with a maximum transmission window size limit.

Motivated by this, in this paper we analyze the effectiveness of this interleaved DBA scheme in a queueing-theoretical context. To this end, we quantify fairness and efficiency measures related to the expected delay figures at each of the queues, providing new fundamental insight in the trade-off between fairness and efficiency by implementing one-stage and two-stage gated service policies. The two-stage gated service policy was introduced in Park et al. [10] where the authors study a symmetric version of the model described in this paper.

A polling system is a multi-queue single-server system in which the server visits the queues in cyclic order to process requests pending at the queues. Polling models occur naturally in the modeling of systems in which service capacity (e.g., CPU, bandwidth) is shared by different types of users, each type having specific traffic characteristics and performance requirements. Polling models find many applications in the areas of computer-communication networks, production, manufacturing and maintenance [9]. Since the late 1960s polling models have received much attention in the literature [12,13]. There are several good reasons for considering heavy-traffic asymptotics, which have recently started to gain momentum in the literature, initiated by the pioneering work of Coffman et al. [2,3] in the mid 90s. Exact analysis of the delay in polling models is only possible in some cases, and even in those cases numerical techniques are usually required to obtain the expected delay at each of the queues. However, the use of numerical techniques for the analysis of polling models has several drawbacks. First, numerical techniques do not reveal explicitly how the system performance depends on the system parameters and can therefore contribute to the understanding of the system behavior only to a limited extent. Exact closed-form expressions provide much more insight into the dependence of the performance measures on the system parameters. Second, the efficiency of each of the numerical algorithms degrades significantly for heavily loaded, highly asymmetric systems with a large number of queues, while the proper operation of the system

is particularly critical when the system is heavily loaded. These observations raise the importance of an exact asymptotic analysis of the delay in polling models in heavy traffic.

We consider an asymmetric cyclic polling model with generally distributed service times and switch-over times. Each queue receives so-called two-stage gated service, which works as follows: Newly incoming customers are first queued at the stage-1 buffer. When the server arrives at a queue, it closes a gate behind the customers residing in the stage-1 buffer, then serves all customers waiting in the stage-2 buffer on a FCFS basis, and moves all customers before the gate at the stage-1 buffer to the stage-2 buffer before moving to the next queue. We focus on the expected delay incurred at each of the queues. First, we derive a so-called pseudo-conservation law for the model under consideration, giving a closed-form expression for a specific weighted sum of the mean waiting times. Then, we describe how the individual expected delays at each of the queues can be calculated recursively by means of the so-called Descendant Set Approach (DSA), proposed for the case of standard one-stage gated and exhaustive service in [5]. The two-stage gated model introduces several interesting complications that do not occur for the standard gated/exhaustive case. Denoting by $X_i^{(k)}$ the number of customers at Q_i in stage k at an arbitrary polling instant at Q_i ($k = 1, 2$), the mean delay at Q_i does not only depend on the first two marginal moments of $X_i^{(2)}$, but also depends on the *cross*-moments $E[X_i^{(1)}X_i^{(2)}]$ (see (7)). To this end, we need to consider the numbers of customers at a queue in both stage 1 and stage 2 at polling instants at that queue, which leads to a two-dimensional analysis.

The main result of this paper is the presentation of a closed-form expression for $(1 - \rho)E[W_i]$, referred to as the scaled expected delay at Q_i , when the load tends to 1. The expression is strikingly simple and shows explicitly how the expected delay depends on the system parameters, thereby explicitly quantifying the trade-off between the increase in fairness and decrease of efficiency introduced by implementing two-stage gated service policies. In particular, the results provide new fundamental insight with respect to mean waiting times for one-stage versus two-stage gated service policies. Furthermore, the results reveal a variety of asymptotic insensitivity properties, which provide new insights into the behavior of polling system under heavy load. The validity of these properties is illustrated by numerical examples. In addition, the expressions obtained suggest simple and fast approximations for the mean delay at each of the queues in stable polling systems. The accuracy of the approximations is evaluated by numerical experiments. The results show that the approximations are highly accurate when the system load is significant.

The remainder of this paper is organized as follows. In section 2 the model is described. In section 3 we present the pseudo-conservation law for the model under consideration. In section 4 we describe how the expected delay figures can be obtained by the use of the DSA. In section 5 we present closed-form expressions for the scaled expected delay at each of the queues, and discuss several asymptotic properties. Finally, in section 6 we propose and test simple

approximations for the moments of the waiting times in heavy traffic, and address the practicality of the asymptotic results.

2 Model Description

Consider a system consisting of $N \geq 2$ stations $Q_1, \dots, Q_N$, each consisting of a stage-1 buffer and a stage-2 buffer. A single server visits and serves the queues in cyclic order. Type- i customers arrive at Q_i according to a Poisson arrival process with rate λ_i , and enter the stage-1 buffer. The total arrival rate is denoted by $\Lambda = \sum_{i=1}^N \lambda_i$. The service time of a type- i customer is a random variable B_i , with Laplace-Stieltjes Transform (LST) $B_i^*(\cdot)$ and with finite k -th moment $b_i^{(k)}$, $k = 1, 2$. The k -th moment of the service time of an arbitrary customer is denoted by $b^{(k)} = \sum_{i=1}^N \lambda_i b_i^{(k)} / \Lambda$, $k = 1, 2$. The load offered to Q_i is $\rho_i = \lambda_i b_i^{(1)}$, and the total offered load is equal to $\rho = \sum_{i=1}^N \rho_i$. Define a polling instant at Q_i as a time epoch at which the server visits Q_i . Each queue is served according to the two-stage gated service policy, which works as follows. When the server arrives at a queue, it closes the gate behind the customers residing in the stage-1 buffer. Then, all customers waiting in the stage-2 buffer are served on a First-Come-First-Served (FCFS) basis. Subsequently, all customers before the gate at the stage-1 buffer are instantaneously forwarded to the stage-2 buffer, and the server proceeds to the next queue. Upon departure from Q_i the server immediately proceeds to Q_{i+1} , incurring a switch-over time R_i , with LST $R_i^*(\cdot)$ and finite k -th moment $r_i^{(k)}$, $k = 1, 2$. Denote by $r := \sum_{i=1}^N r_i^{(1)} > 0$ the expected total switch-over time per cycle of the server along the queues. All interarrival times, service times and switch-over times are assumed to be mutually independent and independent of the state of the system. Necessary and sufficient condition for the stability of the system is $\rho < 1$ (cf. [4]).

Let W_i be the delay incurred by an arbitrary customer at Q_i , defined as the time between the arrival of a customer at a station and the moment at which it starts to receive service. Our main interest is in the behavior of $E[W_i]$. It will be shown that, for $i = 1, \dots, N$,

$$E[W_i] = \frac{\omega_i}{1 - \rho} + o((1 - \rho)^{-1}), \quad \rho \uparrow 1. \quad (1)$$

where the limit is taken such that the arrival rates are increased, while keeping both the ratios between the arrival rates and the service-time distributions fixed. The main result of the paper is a closed-form expression for ω_i , in a general parameter setting (see section 5). Throughout, the following notation is used. For each variable x that is a function of ρ , we denote its values *evaluated at* $\rho = 1$ by $\hat{x}$.

3 Pseudo-conservation Law

In this section we present a pseudo-conservation law (PCL) for the model described above. On the basis of the principle of work decomposition, Boxma and Groenendijk [1] show the following result: For $\rho < 1$,

$$\sum_{i=1}^N \rho_i E[W_i] = \rho \frac{\rho}{1-\rho} \frac{b^{(2)}}{2b^{(1)}} + \rho \frac{r^{(2)}}{2r} + \frac{r}{2(1-\rho)} \left[\rho^2 - \sum_{i=1}^N \rho_i^2 \right] + \sum_{i=1}^N E[M_i], \quad (2)$$

where M_i stands for the amount of work at Q_i at an arbitrary moment at which the server departs from Q_i . Then

$$M_i = M_i^{(1)} + M_i^{(2)}, \quad (3)$$

where $M_i^{(k)}$ is the amount of work at stage k at a server departure epoch from Q_i , $k = 1, 2$. Simple balancing arguments lead to the following expression for $E[M_i]$: For $i = 1, \dots, N$,

$$E[M_i] = E[M_i^{(1)}] + E[M_i^{(2)}] = \rho_i^2 \frac{r}{1-\rho} + \rho_i \frac{r}{1-\rho}. \quad (4)$$

4 The Descendant Set Approach

For $i=1$, define the two-dimensional random variable $\underline{X}_i := (X_i^{(1)}, X_i^{(2)})$, where $X_i^{(k)}$ is the number of stage- k customers at Q_i at an arbitrary polling instant at Q_i when the system is in steady state ($k = 1, 2$), and denote the corresponding Probability Generating Function (PGF) by

$$X_i^*(z_1, z_2) := E \left[z_1^{X_i^{(1)}} z_2^{X_i^{(2)}} \right]. \quad (5)$$

Denoting the Laplace-Stieltjes Transform (LST) of the waiting-time distribution at Q_i by $W_i^*(\cdot)$, the waiting-time distribution at Q_i is related to the distribution of $\underline{X}_i$ through the following expressions (cf. [10]): For $\text{Re } s \geq 0$, $i = 1, \dots, N$, $\rho < 1$,

$$W_i^*(s) = \frac{X_i^*(1 - s/\lambda_i, B_i^*(s)) - X_i^*(1 - s/\lambda_i, 1 - s/\lambda_i)}{E[X_i^{(2)}] (B_i^*(s) - 1 + s/\lambda_i)}. \quad (6)$$

Then it is easy to verify that $E[W_i]$ can be expressed in terms of the first two (cross-)moments of $\underline{X}_i$ as follows: for $i = 1, \dots, N$, $\rho < 1$,

$$E[W_i] = \frac{1}{\lambda_i E[X_i^{(2)}]} \left(\frac{1 + \rho_i}{2} E[X_i^{(2)}(X_i^{(2)} - 1)] + E[X_i^{(1)} X_i^{(2)}] \right). \quad (7)$$

Note that the first part of (7) is similar to the formula for the mean delay in the model with one-stage gated service policy. For this model we have (see Takagi [11])

$$E[W_i^{(one-stage)}] = (1 + \rho_i) \frac{E[X_i(X_i - 1)]}{2\lambda_i E[X_i]}, \quad (8)$$

with X_i the steady-state number of customers at Q_i at an arbitrary polling instant at Q_i .

To derive $E[W_i]$, it is sufficient to obtain the first two factorial moments of $X_i^{(2)}$, and the cross-moments $E[X_i^{(1)}X_i^{(2)}]$. To this end, note first that straight-forward balancing arguments indicate that the first moments $E[X_i^{(k)}]$ ($k = 1, 2$) can be expressed in following closed form:

$$E[X_i^{(1)}] = E[X_i^{(2)}] = \frac{\lambda_i r}{1 - \rho}. \quad (9)$$

However, in general the second-order moments can not be obtained explicitly. In the literature, there are several (numerical) techniques to obtain the moments of the delay. In this section we focus on the Descendant Set Approach (DSA), an iterative technique based on the concept of so-called descendant sets [5]. The use of the DSA for the present model is discussed below.

The customers in a polling system can be classified as originators and non-originators. An originator is a customer that arrives at the system during a switch-over period. A non-originator is a customer that arrives at the system during the service of another customer. For a customer C , define the children set to be the set of customers arriving during the service of C ; the descendant set of C is recursively defined to consist of C , its children and the descendants of its children. The DSA is focused on the determination of the moments of the delay at a fixed queue, say Q_1 . To this end, the DSA concentrates on the determination of the distribution of the two-dimensional stochastic vector $\underline{X}_1(P^*) := (X_1^{(1)}(P^*), X_1^{(2)}(P^*))$, where $X_1^{(k)}(P^*)$ is defined as the number of stage- k customers at Q_1 present at an arbitrary fixed polling instant P^* at Q_1 ($k = 1, 2$). P^* is referred to as the reference point at Q_1 . The main ideas are the observations that (1) each of the customers present at Q_1 at the reference point P^* (either at stage-1 or stage-2) belongs to the descendant set of exactly one originator, and (2) the evolutions of the descendant sets of different originators are stochastically independent. Therefore, the DSA concentrates on an arbitrary tagged customer which arrived at Q_i in the past and on calculating the number of type-1 descendants it has at both stages at P^* . Summing up these numbers over all past originators yields $\underline{X}_1(P^*)$, and hence $\underline{X}_1$, because P^* is chosen arbitrarily.

The DSA considers the Markov process embedded at the polling instants of the system. To this end, we number the successive polling instants as follows. Let $P_{N,0}$ be the last polling instant at Q_N prior to P^* , and for $i = N - 1, \dots, 1$, let $P_{i,0}$ be recursively defined as the last polling instant at Q_i prior to $P_{i+1,0}$. In addition, for $c = 1, 2, \dots$, we define $P_{i,c}$ to be the last polling instant at Q_i prior to $P_{i,c-1}$, $i = 1, \dots, N$. The DSA is oriented towards the determination of the contribution to $\underline{X}_1(P^*)$ of an arbitrary customer present at Q_i at $P_{i,c}$. To this end, define an (i, c) -customer to be a customer present at Q_i at $P_{i,c}$. Moreover, for a tagged (i, c) -customer $T_{i,c}$ at stage 1, we define $\underline{A}_{i,c} := (A_{i,c}^{(1)}, A_{i,c}^{(2)})$, where $A_{i,c}^{(k)}$ is the number of type-1 descendants it has at stage k at P^* , $k = 1, 2$. In this way, the two-dimensional random variable $\underline{A}_{i,c}$ can be viewed as the contribution of $T_{i,c}$ to $\underline{X}_1(P^*)$. Denote the joint PGF of $\underline{A}_{i,c}$ by

$$A_{i,c}^*(z_1, z_2) := E \left[z_1^{A_{i,c}^{(1)}} z_2^{A_{i,c}^{(2)}} \right]. \quad (10)$$

To express the distribution of $\underline{X}_1$ in terms of the distributions of the DS variables $\underline{A}_{i,c}$, denote by $R_{i,c}$ the switch-over period from Q_i to Q_{i+1} immediately after the service period at Q_i starting at $P_{i,c}$. Moreover, denote $\underline{S}_{i,c} := (S_{i,c}^{(1)}, S_{i,c}^{(2)})$, where $S_{i,c}^{(k)}$ is the total contribution to $X_1^{(k)}$ of all customers that arrive at the system during $R_{i,c}$ (note that, by definition, these customers are original customers), and denote the joint PGF of $\underline{S}_{i,c}$ by

$$S_{i,c}^*(z_1, z_2) := E \left[z_1^{S_{i,c}^{(1)}} z_2^{S_{i,c}^{(2)}} \right]. \quad (11)$$

In this way, $\underline{S}_{i,c} = (S_{i,c}^{(1)}, S_{i,c}^{(2)})$ can be seen as the joint contribution of $R_{i,c}$ to $\underline{X}_1$. It is readily verified that we can write

$$\underline{X}_1 = (X_1^{(1)}, X_1^{(2)}) = \sum_{i=1}^N \sum_{c=0}^{\infty} (S_{i,c}^{(1)}, S_{i,c}^{(2)}) = \sum_{i=1}^N \sum_{c=0}^{\infty} \underline{S}_{i,c}. \quad (12)$$

Note that $S_{i,c}^{(1)}$ and $S_{i',c'}^{(2)}$ are dependent if $(i, c) = (i', c')$ but independent otherwise. Hence we can write, for $|z_1|, |z_2| \leq 1$,

$$X_1^*(z_1, z_2) = \prod_{i=1}^N \prod_{c=0}^{\infty} S_{i,c}^*(z_1, z_2). \quad (13)$$

Because $\underline{S}_{i,c}$ is the total joint contribution to $\underline{X}_1$ of all (original) customers that arrive during $R_{i,c}$, the joint distribution of $\underline{S}_{i,c}$ can be expressed in terms of the distributions of the DS-variables $\underline{A}_{i,c}$ as follows: For $i = 1, \dots, N$, $c = 0, 1, \dots$, and $|z_1|, |z_2| \leq 1$,

$$S_{i,c}^*(z_1, z_2) = R_i^* \left(\sum_{j=i+1}^N [\lambda_j - \lambda_j A_{j,c}^*(z_1, z_2)] + \sum_{j=1}^i [\lambda_j - \lambda_j A_{j,c-1}^*(z_1, z_2)] \right). \quad (14)$$

To define a recursion for the evolution of the descendant set, note that a customer at stage-1 present at Q_1 at the polling instant at Q_1 during cycle c is served during the *next* cycle, which lead to the following relation: For $i = 1, \dots, N$, $c = 0, 1, \dots$, and $|z_1|, |z_2| \leq 1$,

$$A_{i,c}^*(z_1, z_2) = B_i^* \left(\sum_{j=i+1}^N [\lambda_j - \lambda_j A_{j,c-1}^*(z_1, z_2)] + \sum_{j=1}^i [\lambda_j - \lambda_j A_{j,c-2}^*(z_1, z_2)] \right), \quad (15)$$

supplemented with the basis for the recursion

$$A_{i,-1}^*(z_1, z_2) = z_1 I_{\{i=1\}}, \quad \text{and} \quad A_{i,-2}^*(z_1, z_2) = z_2 I_{\{i=1\}}, \quad (16)$$

where I_E is the indicator function of the event E . In this way, relations (13)-(16) give a complete, characterization of the distribution of $\underline{X}_1$. Similarly, recursive relations to calculate the (cross-)moments of $\underline{X}_1$ can be readily obtained from those equations.

5 Results

In this section we will present heavy-traffic results that can be proven by exploring the use of the DSA. For compactness of the presentation, the details of the proofs are omitted.

Theorem 1. For $i = 1, \dots, N$,

$$\lim_{\rho \uparrow 1} (1-\rho)^2 E \left[X_i^{(2)} (X_i^{(2)} - 1) \right] = \lim_{\rho \uparrow 1} (1-\rho)^2 E \left[X_i^{(1)} X_i^{(2)} \right] = \hat{\lambda}_i^2 \left[r^2 + \frac{r}{\delta} \frac{b^{(2)}}{b^{(1)}} \right], \quad (17)$$

where

$$\delta := \sum_{i=1}^N \hat{\rho}_i (3 + \hat{\rho}_i). \quad (18)$$

Proof. The result can be obtained along the lines similar to the derivation of the results for the one-stage polling models in [14,15]. $\square$

Theorem 2 (Main result). For $i = 1, \dots, N$,

$$\omega_i = \frac{(3 + \hat{\rho}_i)}{\sum_{j=1}^N \hat{\rho}_j (3 + \hat{\rho}_j)} \frac{b^{(2)}}{2b^{(1)}} + \frac{r(3 + \hat{\rho}_i)}{2}. \quad (19)$$

Proof. The result follows directly by combining (7), (9) and Theorem 1. $\square$

Theorem 2 reveals a variety of properties on the dependence of the limit of the scaled mean waiting times with respect to the system parameters.

Corollary 1 (Insensitivity). For $i = 1, \dots, N$,

- (1) ω_i is independent of the visit order,
- (2) ω_i depends on the switch-over time distributions only through r , i.e., the total expected switch-over time per cycle,
- (3) ω_i depends on the second moments of the service-time distributions only through $b^{(2)}$, i.e., the second moment of the service time of an arbitrary customer.

Corollary 1 is known to be not generally valid for stable systems (i.e., for $\rho < 1$), where the visit order, the second moments of the switch-over times and the individual second moments of the service-time distributions *do* have an impact on the mean waiting times. Hence, Corollary 1 shows that the influence of these

parameters on the mean waiting times vanishes when the load tends to unity, and as such can be viewed as lower-order effects in heavy traffic.

Let us now discuss the trade-off between efficiency and fairness for the one-stage and two-stage gated service policies, using the exact asymptotic results presented in Theorem 2. To this end, denote by $\omega_i^{(one-stage)}$ and $\omega_i^{(two-stage)}$ the heavy-traffic residues of the mean waiting times for the case of one-stage and two-stage gated service at all queues, respectively, defined in (1). For the case of one-stage gated service at all queues, the following results holds (cf. [14]): For $i = 1, \dots, N$,

$$\omega_i^{(one-stage)} = \frac{(1 + \hat{\rho}_i)}{\sum_{j=1}^N \hat{\rho}_j (1 + \hat{\rho}_j)} \frac{b^{(2)}}{2b^{(1)}} + \frac{r(1 + \hat{\rho}_i)}{2}. \quad (20)$$

Also, for the one-stage gated polling model the pseudo-conservation law is given by equation (2), supplemented with

$$E[M_i] = \rho_i^2 \frac{r}{1 - \rho}. \quad (21)$$

Denote by V the amount of waiting work in the system. Then using Little's Law and straightforward arguments it is readily verified that, for $\rho < 1$,

$$E[V] = \sum_{i=1}^N \rho_i E[W_i]. \quad (22)$$

Throughout, we will use $E[V]$ to be the measure of efficiency: for a given set of parameters, a combination of service policies (at each of the queues) is said to be more efficient than another combination of policies if the resulting value of $E[V]$ is smaller. Denoting by $V^{(one-stage)}$ and $V^{(two-stage)}$ the amount of work in the one-stage and two-stage gated model, respectively, the following result follows directly from (2)-(4), (21) and (22).

Corollary 2 (Efficiency). *Two-stage gated service is less efficient than one-stage gated service in the sense that*

$$E[V^{(one-stage)}] < E[V^{(two-stage)}]. \quad (23)$$

Definition of unfairness:

For a given polling model, the unfairness is defined as follows:

$$\mathcal{F} := \max_{i,j=1,\dots,N} \left| \frac{E[W_i]}{E[W_j]} - 1 \right|. \quad (24)$$

Note that according to this definition, the higher $\mathcal{F}$, the *less* fair is the service policy. Note also that the symmetric systems are optimally fair, in the sense that $\mathcal{F} = 0$. The following result follows directly from Theorem 2 and (20).

Corollary 3 (Fairness). *Two-stage gated service is asymptotically more fair than one-stage gated service in the following sense: For $i, j = 1, \dots, N$,*

$$\left| \frac{\omega_i^{(two-stage)}}{\omega_j^{(two-stage)}} - 1 \right| = \left| \frac{3 + \hat{\rho}_i}{3 + \hat{\rho}_j} - 1 \right| < \left| \frac{1 + \hat{\rho}_i}{1 + \hat{\rho}_j} - 1 \right| = \left| \frac{\omega_i^{(one-stage)}}{\omega_j^{(one-stage)}} - 1 \right|. \quad (25)$$

Corollaries 2 and 3 give asymptotic results on the relative fairness and efficiency between one-stage and two-stage gated service. To assess whether similar results also hold for stable systems (i.e., with $\rho < 1$) we have performed extensive numerical validation. The results are outlined below. Consider the asymmetric model with the following system parameters: $N = 2$; the service times at both queues are deterministic with mean $b_1^{(1)} = 0.8$ and $b_2^{(1)} = 0.2$; the switch-over times are deterministic with $r_1^{(1)} = r_2^{(1)} = 1$, and the arrival rates at both queues are equal. For this model, we have calculated the expected waiting times at both queues and the unfairness measure (24), for different values of the load, both for one-stage and two-stage gated service at all queues. Table 1 below shows the results.

Table 1. Mean waiting times and fairness

ρ	one-stage gated			two-stage gated		
	$E[W_1]$	$E[W_2]$	$\mathcal{F}$	$E[W_1]$	$E[W_2]$	$\mathcal{F}$
0.50	3.159	2.465	0.28	7.158	6.468	0.11
0.60	4.241	3.186	0.33	9.239	8.196	0.13
0.70	6.045	4.386	0.38	12.705	11.080	0.15
0.80	9.653	6.788	0.42	19.633	16.868	0.16
0.90	20.475	13.998	0.46	40.400	34.299	0.18
0.95	42.119	28.424	0.48	81.919	69.225	0.18
0.98	107.048	71.708	0.49	206.456	174.075	0.19
0.99	215.262	143.851	0.50	413.997	348.837	0.19

Note that for the model analyzed in Table 1 we have $\hat{\rho}_1 = 4/5$, $\hat{\rho}_2 = 1/5$, so that it is readily seen that for the one-stage gated model $\mathcal{F}$ tends to $1/2$ as ρ goes to 1, whereas for two-stage gated $\mathcal{F}$ tends to $3/16 = 0.1875$ in the limiting case. The results shown in Table 1 show that the two-stage gated service is indeed more “fair” than the one-stage gated service for all values of the load, which suggests that Corollary 3 is also applicable to stable systems. We suspect that this type of results may be proven rigorously; however, a more detailed analysis of the relative fairness is beyond the scope of this paper.

6 Approximation

Equation (1) and Theorem 2 suggest the following approximation for $E[W_i]$ in stable polling systems: For $i = 1, \dots, N$, $\rho < 1$,

$$E[W_i^{(app)}] := \frac{\omega_i}{1 - \rho}, \quad (26)$$

where ω_i is given by Theorem 2. To assess the accuracy of the approximation in (26), in terms of “How high should the load be for the approximation to be accurate?”, we have performed numerical experiments to test the accuracy of the approximations for different values of the load of the system. The relative error of the approximation of $E[W_i]$ is defined as follows: For $i = 1, \dots, N$,

$$\Delta\% := \text{abs} \left(\frac{E[W_i^{(app)}] - E[W_i]}{E[W_i]} \right) \times 100\%. \quad (27)$$

For the model considered in Table 1 above, Table 2 shows the exact (obtained via the DSA discussed in section 3) and approximated values (obtained via (26)) of $E[W_1]$ for different values of the load, both for the model in which all queues receive one-stage gated service, and for the model in which all queues receive two-stage gated service. The results in Table 2 demonstrate that the

Table 2. Exact and approximated values for $E[W_1]$

ρ	one-stage gated			two-stage gated		
	$E[W_1^{(app)}]$	$E[W_1]$	$\Delta\%$	$E[W_1^{(app)}]$	$E[W_1]$	$\Delta\%$
0.50	4.329	3.159	37.04	8.302	7.158	15.98
0.60	5.411	4.241	27.59	10.378	9.239	12.33
0.70	7.214	6.045	19.34	13.837	12.705	8.91
0.80	10.821	9.653	12.10	20.755	19.633	5.71
0.90	21.643	20.475	5.70	41.510	40.400	2.75
0.95	43.286	42.119	2.77	83.020	83.020	1.34
0.98	108.215	107.048	1.09	207.550	206.456	0.53
0.99	216.429	215.262	0.54	415.100	413.997	0.27

relative error of the approximations indeed tends to zero as the load tends to 1, as expected on the basis of Theorem 2. Moreover, the results show that the approximation converges to the limit rather quickly when $\rho \uparrow 1$. Roughly, the results are accurate when the load is 80% or more, which demonstrates the applicability of the asymptotic results for practical heavy-traffic scenarios.

Acknowledgment. The authors wish to thank Onno Boxma for interesting discussions on this topic.

References

1. Boxma, O.J., Groenendijk, W.P.: Pseudo-conservation laws in cyclic-service systems. *J. Appl. Prob.* 24, 949–964 (1987)
2. Coffman, E.G., Puhalskii, A.A., Reiman, M.I.: Polling systems with zero switch-over times: a heavy-traffic principle. *Ann. Appl. Prob.* 5, 681–719 (1995)
3. Coffman, E.G., Puhalskii, A.A., Reiman, M.I.: Polling systems in heavy-traffic: a Bessel process limit. *Math. Oper. Res.* 23, 257–304 (1998)

4. Fricker, C., Jaïbi, M.R.: Monotonicity and stability of periodic polling models. *Queueing Systems* 15, 211–238 (1994)
5. Konheim, A.G., Levy, H., Srinivasan, M.M.: Descendant set: an efficient approach for the analysis of polling systems. *IEEE Trans. Commun.* 42, 1245–1253 (1994)
6. Kramer, G., Muckerjee, B., Pesavento, G.: Ethernet PON: design and analysis of an optical access network. *Phot. Net. Commun.* 3, 307–319 (2001)
7. Kramer, G., Muckerjee, B., Pesavento, G.: Interleaved polling with adaptive cycle time (IPACT): a dynamic bandwidth allocation scheme in an optical access network. *Phot. Net. Commun.* 4, 89–107 (2002)
8. Kramer, G., Muckerjee, B., Pesavento, G.: Supporting differentiated classes of services in Ethernet passive optical networks. *J. Opt. Netw.* 1, 280–290 (2002)
9. Levy, H., Sidi, M.: Polling models: applications, modeling and optimization. *IEEE Trans. Commun.* 38, 1750–1760 (1991)
10. Park, C.G., Han, D.H., Kim, B., Jung, H.-S.: Queueing analysis of symmetric polling algorithm for DBA schemes in an EPON. In: Choi, B.D., (ed.) *Proceedings 1st Korea-Netherlands Joint Conference on Queueing Theory and its Applications to Telecommunication Systems*, Seoul, Korea, pp. 147–154 (2005)
11. Takagi, H.: *Analysis of Polling Systems*. MIT Press, Cambridge, MA (1986)
12. Takagi, H.: Queueing analysis of polling models: an update. In: Takagi, H., (ed.) *Stochastic Analysis of Computer and Communication Systems*, North-Holland, Amsterdam, pp. 267–318 (1990)
13. Takagi, H.: Queueing analysis of polling models: progress in 1990-1994. In: Dshalalow, J.H. (ed.) *Frontiers in Queueing: Models, Methods and Problems*, pp. 119–146. CRC Press, Boca Raton, FL (1997)
14. van der Mei, R.D., Levy, H.: Expected delay in polling systems in heavy traffic. *Adv. Appl. Prob.* 30, 586–602 (1998)
15. van der Mei, R.D.: Polling systems in heavy traffic: higher moments of the delay. *Queueing Systems* 31, 265–294 (1999)
16. van der Mei, R.D.: Delay in polling systems with large switch-over times. *J. Appl. Prob.* 36, 232–243 (1999)