

Approximate Analysis of Central Server Models

Abstract: Service time distributions at computer processing units are often nonexponential. Empirical studies show that different programs may have markedly different processing time requirements. When queuing disciplines are first come, first served, preemptive priority or nonpreemptive priority, models reflecting these characteristics are difficult to analyze exactly. Available approximate techniques are often too expensive for parametric analysis. Inexpensive approximate techniques for solution of central server models with the above characteristics are presented. The results of these techniques are validated with simulation results.

1. Introduction

Central server queuing network models have been widely used in the analysis of computing systems [1-5]. In these models it is assumed that a fixed number of customers (programs) traverse a closed network consisting of the central processor (CPU) and the input/output (I/O) devices. A customer alternately receives service from the CPU and one of the I/O devices. A customer may have to wait in a queue if the server is busy. After completing service at the CPU, a customer selects an I/O device according to probabilities associated with that device and the given customer. These probabilities are independent of the state of the system. The service time of a customer on a device may depend upon the device, the customer, and the queue lengths for that device, but is otherwise independent of the state of the system. Figure 1 illustrates a central server model with three I/O devices.

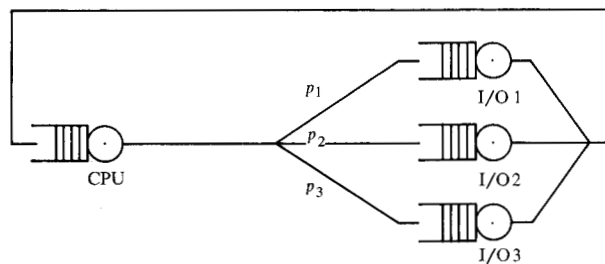
Often the models used are such that solutions for the equilibrium behavior can be determined using the techniques of local balance [6, 7]. If the model is to have first come, first served (FCFS) queuing disciplines, and if the techniques of local balance are to be used in the solution of the model, then it must be assumed that, at the servers with FCFS disciplines, the service distributions are exponential and independent of the customer being served. Local balance techniques do not allow priority queuing disciplines.

Empirical studies on real computing systems show that CPU service distributions are often hyperexponential (the standard deviation is greater than the mean) and that I/O device service distributions may be hypoexponential (the standard deviation is less than the mean). Studies [8] have shown that mean service times and service distributions are dependent on the customer

being served. When one makes assumptions that distributions are exponential and all customers have the same distributions, significant inaccuracy may be introduced into the model. Clearly, distinctions must be made between customers if priority CPU distributions are considered. Therefore, 1) many realistic problems do not satisfy local balance, and 2) customer differentiation is often required for realistic models.

Chandy, Herzog, and Woo [9] have developed very good approximate iterative techniques for analysis of general queuing networks with nonexponential service distributions and distributions dependent on customer class. The iterative techniques of Wallace and Rosenberg [10] may also be used to obtain exact solutions for models with nonexponential distributions. The techniques of Crane and Iglehart [11, 12] may be used to obtain confidence intervals for simulation results for these models and thus to obtain accurate simulation results. However, these techniques are relatively expensive to apply. In many instances it will not be practical to survey a large variety of models using these techniques.

Figure 1 Central server model.



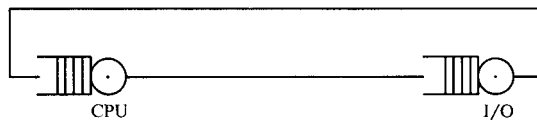


Figure 2 CPU queue and composite I/O queue.

We present here approximate solution techniques specifically intended for, but not limited to, central server models of computing systems. Our techniques are considerably less expensive to apply than the above mentioned techniques, but are sufficiently accurate for the initial stages of computer system design. Our techniques complement the previous techniques in that ours can be used to study and compare a large variety of models, and then more accurate, more expensive techniques may be used to study more carefully a small subset of the original group of models.

In Section 2 of the paper we summarize central server models in local balance and give examples of inaccuracies of "local balance assumptions." In Section 3 we describe "Norton's theorem" on locally balanced queuing networks [13] as applied to central server models. Our approximations are based on the results of Norton's theorem. In Section 4 we present the approximations for models with nonexponential distributions, in Section 5 we present techniques for class dependent service distributions, and in Section 6 we present techniques for models with priority CPU disciplines based on customer class. In Section 7 we compare the results of our techniques with results of simulations; our techniques are validated by comparison with over 125 different simulations.

2. Local balance

A central server model will be in local balance [6] if 1) branching probabilities are dependent only on the device and the customer class, 2) all queuing disciplines are FCFS, processor sharing (PS), or last come, first served preemptive resume (LCFSPR), 3) servers with FCFS discipline have exponential distributions independent of customer class, and 4) servers with PS or LCFSPR disciplines have differentiable service distributions (which may be dependent on customer class). In these models the equilibrium state probabilities will have the "product form" and are easily calculated [6]. From the state probabilities one can determine model statistics such as throughput, server utilization, queue length distributions, and waiting time distributions.

The following example illustrates the inaccuracy which may be introduced by using local balance solutions for models violating local balance assumptions. This example is by no means a worst case, but illustrates

that results of assuming local balance are likely to be unsatisfactory.

Suppose that a system to be modeled has one I/O device and two classes of customers, with one customer per class. Further, both service disciplines are FCFS, all service distributions are exponential, the mean CPU service time for class one is 2, the mean CPU service time for class two is 0.2, and the mean I/O service time for both classes is 1. Suppose we are interested in the overall throughput of customers through the CPU. This model is small enough that exact solution of the Markov balance equations is practical, and from the solution of these equations the throughput is 0.5941. If we assume that the results for a similar model with PS CPU discipline will be accurate enough, the value we get for throughput is 0.84, an error of more than 40 percent. If we apply the techniques of Section 5, the value we get for throughput is 0.6375, an error of about seven percent.

Other examples illustrating the inaccuracy introduced by local balance assumptions are found in [9].

3. Norton's theorem applied to central server models

In this section we review earlier work on Norton's theorem (Subsection 3.1) and present two examples: in Subsection 3.2 we present a multiclass problem and in Subsection 3.3 we work out a single-class example.

• 3.1 Norton's theorem: a discussion

Norton's theorem [13] may be used to transform a central server model in local balance into one with a single "composite" I/O which represents the combined effects of the I/O devices in the original model at steady state. See Fig. 2. Values determined for equilibrium throughputs, server utilizations, and CPU queue length and waiting time distributions of the two-queue model will be the same as those calculated for the original model. The transformation is independent of the CPU parameters, so if a variety of CPU parameters is to be studied, effort may be saved by applying Norton's theorem and studying the reduced model as the CPU parameters are varied. The approximation technique presented here is also especially well suited for parametric analysis of the CPU.

In describing Norton's theorem we shall assume without loss of generality that there are m classes of customers and exactly one customer of each class. The composite I/O processes all customers in parallel in the two-queue, CPU-composite I/O model. The composite I/O service rate for the customer of any given class i at any given time depends upon i and upon the number of customers N_j of class j ($N_j = 0$ or 1), $j = 1, \dots, m$, in the composite I/O queue at that time. These composite

I/O service rates are determined by analyzing a modified version of the original network in which the CPU has been "shorted," i.e., the mean CPU service time for all customers is set to zero. See Fig. 3. The composite I/O service rate of a customer of any given class i , when there are N_j customers of class j ($N_j = 0$ or 1), $j = 1, \dots, m$, in the composite I/O queue, is set equal to the throughput of the customer in class i through the shorted CPU when there is a population of N_j customers of class j , $j = 1, \dots, m$ in the shorted CPU model. The solutions of the two-queue, CPU-composite I/O model, with the same CPU parameters as in the original model and these queue-dependent composite I/O service rates, will be identical to those of the original model for the equilibrium statistics mentioned above.

• 3.2 Example 1

Consider the following two-class example of a locally balanced central-server model with a processor-shared CPU and two I/Os labeled 1 and 2, two nonidentical customers, one of class A and the other of class B. The class A customer uses I/Os 1 and 2 with equal probability, while the class B customer uses I/O 1 exclusively. The mean service time for each I/O is independent of customer class. The mean service times for I/Os 1 and 2 are 1 and 2, respectively. All I/O service times have negative-exponential distributions.

Both class A and B customers are assumed to be serviced in parallel in the composite I/O queue. The service rates for class A and B customers depend upon the numbers of class A and B customers in the composite I/O queue. We next discuss the computation of these rates by analyzing the modified version of the original network in which the CPU has been shorted (Fig. 3). When only the class A customer is present in the CPU-shortened network, the throughput of the class A customer through the shorted CPU is $\frac{2}{3}$; when only the class B customer is present the throughput is 1; and when both are present, the throughputs for classes A and B are $\frac{1}{2}$ and $\frac{3}{4}$, respectively. The composite I/O service rates when there is one customer of class A and none of class B in the composite I/O queue is set to $\frac{2}{3}$ for class A (and 0 for class B); when there is one customer of class B and none of class A the rate is set to 1 for class B (and 0 for class A); and when there is one customer of each class the rates are set to $\frac{1}{2}$ for class A and $\frac{3}{4}$ for class B. The solution of the CPU-composite I/O model with the same CPU parameters as in the original model and these queue-dependent composite I/O service rates will be identical to the solutions of the original model for the equilibrium statistics mentioned above.

• 3.3 Example 2

Consider the following single-customer-class example. A locally balanced central server model has two identical

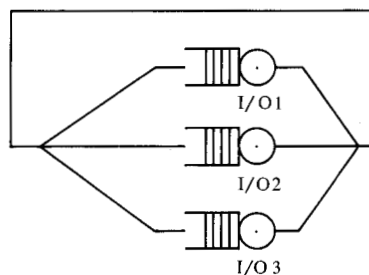


Figure 3 Central server model with shorted CPU.

customers and two I/O devices labeled 1, 2. The probabilities that a customer will branch to the i th I/O device, $i = 1, 2$, are 0.5 and 0.5, respectively. The mean service times for I/Os 1 and 2 are 4.0 and 2.0, respectively. All I/O queues have the FCFS discipline. With the CPU shorted and j customers in the shorted network, $j = 1, 2$, the throughputs through the shorted CPU are $\frac{1}{3}$ and $\frac{3}{7}$, respectively. In the two-queue, CPU-composite I/O model, the total service rates of the composite I/O when there are j customers in the composite I/O queue, $j = 1, 2$, are set to $\frac{1}{3}$ and $\frac{3}{7}$, respectively. Equivalently, with all customers served in parallel, the rate for each customer when there are j customers in the queue is $\frac{1}{3}$ for $j = 1$ and $(\frac{1}{2}) \times (\frac{3}{7}) = \frac{3}{14}$ for $j = 2$.

The equilibrium statistics computed for the CPU-composite I/O model with the same CPU parameters as in the original model, and these queue-length-dependent composite I/O service rates, will be identical to the equilibrium statistics computed for the original model.

4. FCFS central server models with nonexponential service times

We first discuss the overall technique generally (4.1), then study composite I/O representations (4.2), present the detailed algorithm (4.3), and work out an example (4.4).

• 4.1 Overview

We now restrict our attention to central server models with all customers identical, FCFS disciplines at all servers, and arbitrary service distributions having rational Laplace transforms. Even though this class of models is not in local balance except when all service distributions are exponential, we shall apply Norton's theorem and show that the composite I/O model yields solutions close to those of the original model. (In making the composite I/O transformation we assume that the I/O devices have exponential distributions with the same means as the actual distributions. See example below.) Chandy, Herzog, and Woo [9] use an approximate application of Norton's theorem in their iterative method.

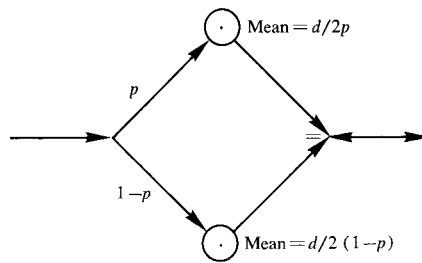


Figure 4 Hyperexponential distribution model.

In order to compensate for the inaccuracy introduced, we adjust the distributions for the composite I/O to reflect the nonexponential character of the actual distributions.

After applying Norton's theorem and adjusting the distributions, we have a central server model with a single composite I/O, with both service distributions nonexponential. This model is solved by an efficient recursive technique which is an application of the technique developed by Herzog, Woo, and Chandy [14]. Their technique assumes distributions of the generalized Erlang form developed by Cox [15]. This generalized form includes arbitrary distributions with rational Laplace transforms. Our technique assumes that both the CPU and the I/O distributions are of this general form. Details of our two-queue analysis are given in [16].

Our adjustment for the nonexponential nature of the I/O distributions is simple and effective. More sophisticated adjustments could potentially increase the accuracy of the final results. We characterize each I/O distribution by its mean and coefficient of variation (standard deviation divided by the mean). For the means of the composite I/O distributions we use the queue-length-dependent values as shown earlier. We assume that the composite I/O coefficient of variation is the weighted sum of the coefficients of variation of the individual distributions, with the weights being the I/O branching probabilities. The composite I/O coefficient of variation is a constant, independent of queue length. Of course, the mean and coefficient of variation do not completely specify the distribution. If the composite I/O coefficient of variation is greater than one, we assume that the composite I/O service time is a standard two-stage hyperexponential as in Fig. 4. If the coefficient of variation is one, we assume the service time is exponential. If the coefficient of variation is less than one, we assume the service time is of the generalized Erlang form with the minimum number of stages necessary to obtain the given coefficient of variation, all stages having the same mean, and all branching probabilities zero, with the possible exception of the branch after the first stage, as in Fig. 5.

• 4.2 Composite I/O distribution

We desire the composite I/O distribution to represent the aggregate of all the individual I/O distributions. Intuitively, we expect the distribution of a given I/O to influence the composite I/O distribution more than distributions of other I/Os, if the given I/O processes more customers than other I/Os. We decided to restrict attention to the first two moments to keep computation simple. The means of composite I/O service times are obtained by aggregating individual I/O mean service times via Norton's theorem. The composite coefficient of variation is obtained by aggregating individual coefficients of variation, weighting each I/O by its branching probability since I/O branching probabilities are directly proportional to I/O throughputs. Note that although the mean composite service time is queue-length-dependent, the coefficient of variation is not dependent on queue length. Note also that if all the I/Os have the same coefficient of variation, then the composite I/O will have that coefficient of variation as well.

The first two moments do not completely specify a distribution. We decided to model composite service times using either two-stage hyperexponential (Fig. 4) or generalized Erlang (Fig. 5) random variables, since these are common ways of representing service times in computing systems. Note that the particular forms of the hyperexponential and generalized Erlang random variables are such that the first two moments uniquely specify the distributions. The selection of these particular composite I/O distributions was made with modeling convenience and reasonability in mind; clearly other choices could also have been made. However, note that if the original model satisfies local balance, then our technique gives *exact* results, since the composite I/O distribution obtained via our technique is the same as that obtained via Norton's theorem.

Hyperexponential Let k_c be the coefficient of variation of the composite I/O. We shall use a standard hyperexponential random variable to model composite I/O service times if $k_c > 1$. The relationship between k_c and parameter p (Fig. 4) of the hyperexponential is as follows:

$$p = \frac{k_c^2 + 1 - (k_c^4 - 1)^{\frac{1}{2}}}{2(k_c^2 + 1)}. \quad (1)$$

Note that k_c uniquely specifies p . The means for each stage of this hyperexponential are uniquely specified by p and the mean composite service time, d .

$$\text{Mean of stage 1} = d/2p. \quad (2)$$

$$\text{Mean of stage 2} = d/2(1-p). \quad (3)$$

Generalized Erlang Consider the Erlang random variable (Fig. 5) with n stages, $n = 2, 3, 4, \dots$. After a customer completes the first stage, he may finish service with probability p , or he may continue through the remaining $n - 1$ stages with probability $1 - p$. All stages have the same mean time, and all stage holding times are independent exponential random variables. By varying p from 0 to 1 the coefficient of variation ranges from $1/n^2$ to 1. We wish to keep the number of stages small to minimize computation. Hence, we shall use n stages if and only if $1/(n - 1) > k_c \geq 1/n^2$. The value of n is directly determined from k_c . Together n and k_c uniquely specify p . See Eq. (4) below. The means for each stage are uniquely specified by n , k_c , p , and the means of the composite I/O service times.

$$p = \frac{2nk_c^2 + n - 2 - (n^2 + 4 - 4nk_c^2)^{\frac{1}{2}}}{2(k_c^2 + 1)(n - 1)}. \quad (4)$$

$$\text{Mean of each stage} = d/[n - p(n - 1)]. \quad (5)$$

In conclusion, the generalized Erlang and hyperexponential random variables shown in Figs. 4 and 5 are completely specified by the first two moments and have a wide range of coefficients of variation. The parameters p are independent of composite I/O mean service times and the mean times for all stages in both distributions are directly proportional to the composite I/O mean service time; this simple relationship is an advantage in modeling queue-dependent service rates.

• 4.3 Algorithm

We now present the algorithm after explaining some notation. Let there be R I/O queues indexed $1, \dots, r, \dots, R$. We shall use the subscript r to denote the r th I/O in the original model and the subscript c to denote the composite I/O in the CPU-composite I/O model. Let p_r be the probability that a customer branches to the r th I/O device after finishing CPU service. Let k denote the coefficient of variation: k_c for the composite I/O and k_r for the r th I/O device. We shall use the subscript 0 (zero) for the CPU. Let U_r be the utilization and t_r the throughput for the r th queue, $r = 0, 1, \dots, R$. Let λ_r be the service rate for the r th I/O device. Let $\bar{q}$ and $\bar{w}$ be the mean CPU queue length and wait times and let σ_q and σ_w be the corresponding standard deviations.

Algorithm A

Step 1 Composite I/O service rates. Consider the given (nonlocally balanced) model. Construct the shorted-CPU model in which all I/O service times are assumed to be independent exponential random variables and the CPU service time is set to zero. The shorted-CPU model satisfies local balance and can be analyzed easily.

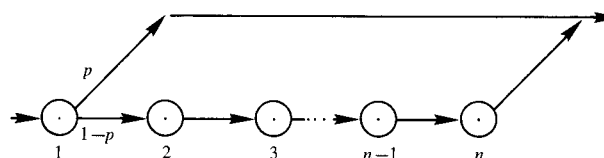


Figure 5 Generalized Erlang distribution model for n stages.

Determine queue-dependent composite I/O service rates by analyzing the shorted-CPU model.

Step 2 Composite I/O coefficient of variation. Compute

$$k_c = \sum_{r=1}^R k_r \cdot p_r.$$

Step 3 Determine exponential stage representations for composite I/O service times from k_c and composite I/O mean service times. If $k_c > 1$, use standard hyperexponential random variable (Fig. 4). If $k_c = 1$, use exponential random variable. If $k_c < 1$, use generalized Erlang random variable (Fig. 5).

Step 4 Solve the two-queue, CPU-composite I/O model. The CPU parameters in this model are set to the same values as in the original model. The composite I/O parameters are completely and uniquely specified by step 3. The two-queue model is completely specified. Analyze this model to determine U_0 , t_0 , $\bar{q}$, σ_q , $\bar{w}$, and σ_w .

Step 5 I/O utilizations. Compute $t_r = t_0 \times p_r$ for $r = 1, \dots, R$, $U_r = t_r/\lambda_r$ for $r = 1, \dots, R$. Stop.

• 4.4 Example

Consider Example 2 (Subsection 3.3) except that I/O 1 has an exponential service time, I/O 2 has a generalized Erlangian service time with a coefficient of variation of 0.414, and the CPU has a standard hyperexponential service time with a coefficient of variation of 2 and a mean of 2. We shall follow through the five steps of the algorithm.

Step 1 The composite I/O service rates (from Subsection 3.3) when there are j customers in the composite I/O queue are $1/3$ and $3/7$, respectively.

Step 2 $k_c = (0.5 \times 1.0) + (0.5 \times 0.414) = 0.707$.

Step 3 Since $k_c < 1$ the generalized Erlang representation is used. In this case n will be 2 and p will be zero. (The rate for each stage is clearly twice the composite I/O service rate.)

Step 4 We now have a two-queue model where the CPU service time is a two-stage hyperexponential and the

composite I/O service time is a two-stage Erlang. The balance equations for the resulting Markov states are solved to obtain $U_0 = 0.571$, $t_0 = 0.286$, $\bar{q}_0 = 0.837$, $\bar{w}_0 = 2.93$.

Step 5 $t_1 = t_0 \times 0.5 = 0.143$, $t_2 = t_0 \times 0.5 = 0.143$, $U_1 = t_1/\lambda_1 = 0.571$, $U_2 = t_2/\lambda_2 = 0.286$. Stop.

5. FCFS central server models with class-dependent service rates

This section is divided into three subsections. In 5.1 we discuss the technique generally, in 5.2 we present the algorithms, and in 5.3 we present an example.

• 5.1 Discussion

In this subsection we restrict ourselves to models with several classes of customers, FCFS, all service distributions exponential, all I/O service rates independent of customer class, and the CPU service rates dependent on customer class [17]. The assumption of class-independent I/O service rates can be justified by observing that the largest portion of most I/O services is spent on primarily program-independent operations such as acquiring channels, positioning disk arms, and waiting for device rotation. The techniques presented here have been extended to nonexponential CPU distributions and can also easily be extended to nonexponential I/O distributions. They are extended to priority disciplines in the next section, using the techniques of the current section. Our techniques may also be extended to other more general models.

When we consider such central server models, even the reduced model obtained by applying the Norton's theorem approximation to the I/O subnetwork is difficult to analyze. As the number of classes and/or the numbers of customers per class attain even moderate values, e.g., 4, the analysis becomes too complex to be of practical value.

To reduce the complexity of analysis, we transform the more general original model to an approximately equivalent one with only two classes of customers: a designated class with only one customer and a composite class representing all of the other customers in the network. This further reduced model can be analyzed relatively easily, by applying the Norton's theorem approximation. By designating each class in the original model and in turn analyzing the corresponding reduced model, we can obtain approximate values for the interesting statistics for each customer class in the original model.

In transforming the original model to the one with only two classes, the customer of the designated class is given the same I/O branching probabilities and CPU service distribution as in the original model. For each I/O device, the composite class branching probability is

determined as a weighted sum of the branching probabilities of the classes being "coalesced" from the original model. The weights used are the relative throughputs of the corresponding customers in a model identical to the original model, except that the CPU is processor-shared; this PS model satisfies local balance and is easily analyzed. The CPU service distribution for the composite class is chosen to be the standard two-stage hyperexponential distribution with mean and second moment determined from weighted sums of the means and second moments of the CPU service distributions of the classes being coalesced from the original model.

After this transformation is applied, the Norton's theorem approximation is applied. The resulting model, with the composite class and composite I/O queue, is analyzed by techniques similar to those used in Section 4.

• 5.2 Algorithms

In this subsection we describe two algorithms: the main program, algorithm B, and a subprogram, algorithm C, which approximates an N -class problem by a two-class problem.

Algorithm B

Assume that there are N classes of customers. For purposes of exposition, we assume (without loss of generality) that there is only one customer in each class.

Step 1 For each class i in turn, $i = 1, \dots, N$, do steps 2i through 5i and thus compute the throughputs and utilizations for all queues for class i , and also the means and variances of CPU queue lengths and wait times for class i . The algorithm stops after all N classes have been considered.

Step 2i Use algorithm C to approximate the given N -class problem by a two-class problem where the two classes are the designated class and a "coalesced class" which represents all customers except those in the designated class. We refer to the original central server model as model A and this two-class approximation as model B. Note that B and A have exactly the same central server network structure; only the number of classes is changed. The parameters for the designated class are the same in A and B. For the coalesced class CPU service time is assumed to be hyperexponential in B. In A and B, I/O service times are identical.

Step 3i Compute composite I/O service rates for the designated and coalesced classes of model B in the usual manner (i.e., by computing throughputs through the shorted CPU of model B and assuming all I/O service times are exponential).

Step 4i Consider the resulting two-queue, two-class network consisting of the CPU and I/O queues and the

designated and coalesced classes; we refer to this network as model C. Solve Markov balance equations to determine steady state probabilities of model C. Determine CPU throughput t_{0i} , utilization U_{0i} , mean and variance of CPU queue length, and wait time for designated class i from the equilibrium state probabilities of model C. Statistics for the coalesced class are not computed.

Step 5i Determine I/O throughputs t_{ri} , and utilizations U_{ri} , for each I/O r , $r = 1, \dots, R$, for the designated class i . Let p_{ri} be the probability that a customer of class i branches to I/O r after CPU service. Then $t_{ri} = t_{0i} \times p_{ri}$ for $r = 1, \dots, R$ and $U_{ri} = t_{ri} / \lambda_i$ for $r = 1, \dots, R$. Statistics for the coalesced class are not computed.

Figure 6 shows the relationships between models A, B, and C.

Algorithm C

In this algorithm we determine CPU service distributions and I/O branching probabilities for the coalesced class.

Step 1 Consider a network identical to the given network (model A) except that the CPU is processor-shared; we shall refer to this network as model D. Model D satisfies local balance and is easily analyzable. For the purposes of algorithm C *only*, we shall approximate the CPU throughputs of model A by those of model D. Compute t'_j , the CPU throughput of class j in model D, for $j = 1, \dots, N$.

Step 2 Compute the conditional probability V_j that a random customer who finishes I/O service in model D is in class j , given that he is not in designated class i :

$$V_j = t'_j / \sum_{h \neq i} t'_h \text{ for } j \neq i;$$

$$= 0 \quad \text{for } j = i.$$

Step 3 Compute the first two moments of the CPU service time for the coalesced class. Let $E[S^n]$ and $E[S_j^n]$ be the n th moments of the CPU service time for the coalesced class and class j , respectively, $j = 1, \dots, N$. Then

$$E[S] = \sum_j E[S_j] \cdot V_j, \quad E[S^2] = \sum_j E[S_j^2] \cdot V_j.$$

Represent CPU service time for the coalesced class by a standard hyperexponential random variable (Fig. 4) with the above first two moments.

Step 4 Approximate I/O branching probabilities for the composite class by

$$p_{rc} = \sum_j p_{rj} \cdot V_j.$$

Stop.

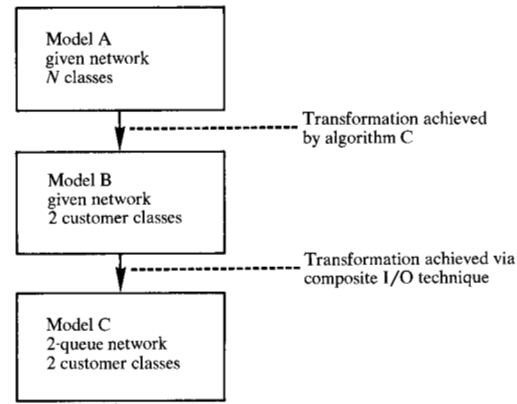


Figure 6 Relationships among models A, B, and C.

5.3 Example

Consider a model with two I/Os and three classes of customers. The mean service times for I/O 1 and I/O 2 are both 2 time units. The branching probabilities for the first I/O are 1, 0, 0.5 for classes 1, 2, and 3, respectively, and 0, 1, 0.5 for the second I/O. CPU mean times for classes 1, 2, 3 are 1, 2, 3, respectively. All service times are assumed to be independent, exponential random variables.

Algorithm B, step 1. We shall carry out steps 2i through 5i for $i = 1$.

We first call algorithm C to obtain the two-class approximation.

Algorithm C, step 1. Analyzing model D we get $t'_2 = 0.159$, $t'_3 = 0.111$.

Algorithm C, step 2. $V_1 = 0$, $V_2 = 0.589$, $V_3 = 0.411$.

Algorithm C, step 3.

$$E[S] = (2 \times 0.589) + (3 \times 0.411) = 2.41,$$

$$E[S^2] = (8 \times 0.589) + (18 \times 0.411) = 12.12.$$

The hyperexponential representation for the CPU service time has parameter $p = 0.398$.

Algorithm C, step 4.

$$p_{1c} = (0 \times 0.589) + (0.5 \times 0.411) = 0.206,$$

$$p_{2c} = (1 \times 0.589) + (0.5 \times 0.411) = 0.795.$$

We now have a two-class problem. The CPU service time for the coalesced class is hyperexponential with mean 2.41; the I/O branching probability for device 1 is 0.206 and for device 2 is 0.795.

Algorithm B, step 3i. The composite I/O service rates for class 1 and the coalesced class for different queue conditions are shown in Table 1.

Algorithm B, step 4i. Model C is analyzed to obtain $t_{01} = 0.173$, $U_{01} = 0.173$, $\bar{q}_1 = 0.59$, $w_1 = 3.43$, $\sigma_{q_1} = 0.49$, $\sigma_{w_1} = 3.09$.

Table 1 Different queue conditions and service rates.

<i>Number of class 1 customers</i>	<i>Number of coalesced class customers</i>	<i>Total rate for class 1</i>	<i>Total rate for coalesced class</i>
1	0	0.5	0
0	1	0	0.5
0	2	0	0.598
1	1	0.415	0.415
1	2	0.386	0.556

Table 2 Hyperexponential CPU, exponential I/O.

<i>Number of customers</i>	<i>2</i>	<i>4</i>	<i>8</i>	<i>12</i>
CPU utilization	0.590	0.713	0.802	0.846
(simulation)	0.588	0.715	0.835	0.864
(local balance)	0.623	0.783	0.884	0.921
CPU mean queue length	0.912	1.85	3.66	5.51
	0.894	1.82	3.79	5.65
CPU standard deviation of queue length	0.851	1.56	2.83	4.26
	0.840	1.54	2.80	4.10
CPU mean wait time	1.54	2.59	4.56	6.52
	1.50	2.57	4.55	6.49
CPU standard deviation of wait time	2.90	4.02	5.71	7.14
	2.73	4.07	5.53	7.24
I/O 1 utilization	0.590	0.713	0.802	0.846
	0.606	0.730	0.824	0.864
I/O 2 utilization	0.148	0.178	0.201	0.212
	0.147	0.176	0.214	0.212
I/O 3 utilization	0.036	0.045	0.050	0.053
	0.036	0.043	0.054	0.056
CPU mean service	1.0			
Coefficient of variation	2.134			
I/O 1 mean service	2.0			
Coefficient of variation	1.0			
Branching probability	0.5			
I/O 2 mean service	1.0			
Coefficient of variation	1.0			
Branching probability	0.25			
I/O 3 mean service	0.25			
Coefficient of variation	1.0			
Branching probability	0.25			

6. Approximations for models with priority CPU disciplines

Now we consider central server models with the same characteristics as in the previous section, except that the CPU discipline will be a priority discipline with priority based on customer class. We will restrict consideration

to preemptive and nonpreemptive priorities based on customer class, but these techniques are directly applicable to other priority disciplines.

Again, we do not try to apply the Norton's theorem approximation directly, but rather coalesce the classes of customers in the original model to simplify the analysis. The reduced model we consider has three classes of customer: a designated class, which we do not restrict to a single customer as in the FCFS model, and two composite classes, one of a higher priority than the designated class and one of lower priority. The coalescing of classes into these three classes is similar to the technique used in the previous section. The coalescing is done separately for each of the two composite classes. The CPU distribution used for each of the composite classes is an exponential distribution with mean taken as the weighted sum of the means of the classes being coalesced into that composite class. The weights are the relative throughputs of classes within the composite class. In other respects, the analysis is essentially the same as that already described.

7. Validation, implementation and performance

We have constructed a simulator which employs the confidence interval techniques of Crane and Iglehart [11, 12]. This simulator can be used with general queuing networks with a variety of disciplines, heterogeneous classes of customers, and generalized Erlang service distributions. The simulator determines confidence intervals during the simulation, and continues the simulation until satisfactory intervals are obtained. Details of the simulator are found in [16]. This simulator has been used to determine results for the various models described below. Crane and Iglehart show how to obtain confidence intervals for results of simulations of Markov models with finite or countable state spaces. In general, the confidence intervals obtained are as follows: For utilization, the 90 percent intervals are at most 0.05 wide. For those cases where queue lengths and waiting times are obtained, the 90 percent intervals for the means are at most ± 6 percent of the point estimates, and the 80 percent intervals for the standard deviations are at most ± 16 percent of the point estimates. In many of the cases the intervals are considerably tighter. However, we were unable to obtain confidence intervals for the FCFS models with six classes of customers. For these models the state space is very large, and we were unable to select a state that the system would return to frequently; this is necessary to apply the Crane and Iglehart techniques. We used predetermined simulation run lengths for the six-class FCFS models, with the run lengths based on experience with four-class FCFS models.

We have implemented our approximation techniques as a set of FORTRAN programs for a CDC 6600. Over 125

Table 3 Exponential CPU, Erlang I/O.

Number of customers	2	4	8	12
CPU utilization	0.646	0.813	0.906	0.937
(simulation)	0.652	0.821	0.914	0.940
(local balance)	0.623	0.783	0.884	0.921
CPU mean queue length	0.881	1.83	3.78	5.77
	0.917	1.90	3.88	5.99
CPU standard deviation of queue length	0.759	1.28	2.39	3.52
	0.778	1.31	2.37	3.62
CPU mean wait time	1.36	2.25	4.18	6.16
	1.43	2.34	4.23	6.42
CPU standard deviation of wait time	1.26	1.83	2.97	4.12
	1.30	1.86	2.92	4.29
I/O 1 utilization	0.646	0.813	0.906	0.937
	0.629	0.801	0.904	0.921
I/O 2 utilization	0.162	0.203	0.227	0.234
	0.165	0.213	0.228	0.232
I/O 3 utilization	0.040	0.051	0.057	0.059
	0.040	0.051	0.057	0.060
CPU mean service	1.0			
Coefficient of variation	1.0			
I/O 1 mean service	2.0			
Coefficient of variation	0.707			
Branching probability	0.5			
I/O 2 mean service	1.0			
Coefficient of variation	0.707			
Branching probability	0.25			
I/O 3 mean service	0.25			
Coefficient of variation	0.707			
Branching probability	0.25			

Table 4 Hyperexponential CPU, Erlang I/O.

Number of customers	2	4	4	12
CPU utilization	0.604	0.730	0.815	0.857
(simulation)	0.604	0.751	0.851	0.893
(local balance)	0.623	0.783	0.884	0.921
CPU mean queue length	0.904	8.81	3.58	5.42
	0.912	1.88	3.75	5.92
CPU standard deviation of queue length	0.828	1.52	2.89	4.21
	0.834	1.50	2.73	3.99
CPU mean wait time	1.50	2.48	4.39	6.33
	1.51	2.57	4.44	6.68
CPU standard deviation of wait time	2.89	4.00	5.67	7.10
	2.86	4.15	5.31	6.89
I/O 1 utilization	0.604	0.730	0.815	0.857
	0.610	0.732	0.845	0.874
I/O 2 utilization	0.151	0.183	0.204	0.214
	0.150	0.187	0.206	0.232
I/O 3 utilization	0.038	0.046	0.051	0.054
	0.038	0.045	0.053	0.056
CPU mean service	1.0			
Coefficient of variation	2.134			
I/O 1 mean service	2.0			
Coefficient of variation	0.707			
Branching probability	0.5			
I/O 2 mean service	1.0			
Coefficient of variation	0.707			
Branching probability	0.25			
I/O 3 mean service	0.25			
Coefficient of variation	0.707			
Branching probability	0.25			

models have been validated to assure a thorough sampling of problems.

Fifty-six of the models validated are of the class described in Section 4, i.e., single-class, nonexponential. These models included from 2 to 12 customers, 1 or 2 CPUs, from 3 to 8 I/O devices, and a wide variety of combinations of distributions, with coefficients of variation ranging from 0.577 to 5.0. In general the models were fairly well balanced, but some of the models were strongly CPU bound or I/O bound. Error tolerances were determined in the manner used in [9] for utilizations, CPU queue lengths, and CPU waiting times. Results are said to be within a tolerance z if 1) the difference in utilization is not more than z , 2) the differences in the means and standard deviations of queue length are not more than z times the number of customers in the network, and 3) the differences in the means and standard deviations of

the wait times are not more than z times the cycle time. For the 56 models studied, the results are generally within a tolerance of 0.05, with a maximum tolerance of 0.18. In [9] a tolerance of 0.05 is considered to be good, and a tolerance of 0.10 is considered adequate. By these standards the results are good or adequate for 51 of the 56 models. For these models, the computer time required per model was negligible, approximately 75 ms per model. Tables 2, 3, and 4 show results for 12 of these models.

Forty-four models of the class described in Section 5, i.e., FCFS with different classes of customers, including four with hyperexponential CPU distributions, have been validated. These models include from 2 to 8 customers, with from 2 to 6 classes of customers, and 3 or 4 I/O devices. Utilizations and throughputs, both overall and by class, were validated for all of these models. For eight

Table 5 FCFS, 4 classes, 1 customer/class, 3 I/O devices.

Class	1	2	3	4	All
CPU throughput	0.182	0.192	0.110	0.168	0.652
(simulation)	0.144	0.156	0.122	0.126	0.548
(local balance)	0.230	0.272	0.090	0.208	0.800
CPU utilization	0.091	0.048	0.552	0.084	0.775
	0.072	0.039	0.612	0.063	0.786
	0.115	0.068	0.451	0.104	0.738
CPU mean queue length	0.36	0.35	0.59	0.35	
(simulation)	0.51	0.49	0.65	0.48	
CPU standard deviation	0.48	0.48	0.49	0.48	
of queue length	0.50	0.50	0.48	0.50	
CPU mean wait time	1.97	1.82	5.32	2.08	
	3.45	3.22	5.44	3.67	
CPU standard deviation	3.95	3.87	4.94	4.08	
of wait time	4.63	4.58	5.13	4.70	
I/O 1 utilization	0.219	0.154	0.044	0.067	
	0.176	0.121	0.047	0.053	
I/O 2 utilization	0.058	0.123	0.071	0.054	
	0.047	0.095	0.077	0.044	
I/O 3 utilization	0.097	0.103	0.118	0.268	
	0.080	0.081	0.128	0.214	
CPU mean service	0.500	0.250	5.00	0.500	
Coefficient of variation	1.00	1.00	1.00	1.00	
I/O 1 probability	0.6	0.4	0.2	0.2	
(mean service = 2.00)					
I/O 2 probability	0.2	0.4	0.4	0.2	
(mean service = 1.60)					
I/O 3 probability	0.2	0.2	0.4	0.6	
(mean service = 2.67)					

Table 6 FCFS, 4 classes, 1 customer/class, 3 I/O devices.

Class	1	2	3	4	All
CPU throughput	0.111	0.182	0.162	0.152	0.607
(simulation)	0.118	0.142	0.130	0.128	0.518
(local balance)	0.093	0.239	0.216	0.230	0.778
CPU utilization	0.555	0.091	0.081	0.038	0.765
	0.590	0.071	0.065	0.032	0.758
	0.465	0.119	0.108	0.058	0.750
CPU mean queue length	0.61	0.39	0.37	0.36	
(simulation)	0.64	0.52	0.50	0.48	
CPU standard deviation	0.49	0.49	0.48	0.48	
of queue length	0.48	0.50	0.50	0.50	
CPU mean wait time	5.51	2.13	2.30	2.37	
	5.28	3.69	3.75	3.68	
CPU standard deviation	5.08	6.90	7.25	7.56	
of wait time	9.48	8.67	8.83	8.89	
I/O 1 utilization	0.133	0.145	0.064	0.061	
	0.145	0.115	0.053	0.050	
I/O 2 utilization	0.036	0.116	0.103	0.049	
	0.040	0.089	0.086	0.041	
I/O 3 utilization	0.059	0.097	0.172	0.245	
	0.062	0.077	0.149	0.209	
CPU mean service	5.00	0.500	0.500	0.250	
Coefficient of variation	2.00	2.00	2.00	2.00	
I/O 1 probability	0.6	0.4	0.2	0.2	
(mean service = 2.00)					
I/O 2 probability	0.2	0.4	0.4	0.2	
(mean service = 1.60)					
I/O 3 probability	0.2	0.2	0.4	0.6	
(mean service = 2.67)					

of the models, queue lengths and wait times for each class were also validated. We did not explicitly determine tolerances as in the single-class models, but in general the results showed good accuracy for utilization and reasonable accuracy overall. Tables 5 and 6 give results for two of the models. For the 44 models, the programs required approximately 400 ms of computation per model.

Thirty-six priority models were validated, 24 preemptive and 12 nonpreemptive. These models included from 4 to 6 customers, with from 3 to 6 classes, and 3 or 4 I/O devices. Again, utilizations and throughputs were validated for all models. Central processing unit queue lengths and mean CPU wait times were validated for 12

preemptive models and all nonpreemptive models. Tables 7 and 8 show results for two models. For the 36 models, the computation per model was approximately 400 ms.

In addition to providing reasonable accuracy for models not in local balance, these programs give exact results for models in local balance where class coalescing is not necessary. Though the coalescing techniques do not necessarily give exact results for locally balanced models, the results are very close. In the above validation process, for all FCFS models requiring coalescing, the coalescing process was applied to a locally balanced model similar to the nonlocally balanced model being studied. Individual class throughputs and utilizations

Table 7 Preemptive, 6 classes, 1 customer/class, 3 I/O devices.

<i>Class</i>	<i>1</i>	<i>2</i>	<i>3</i>	<i>4</i>	<i>5</i>	<i>6</i>	<i>All</i>
CPU throughput	0.119	0.166	0.149	0.138	0.143	0.138	0.873
(simulation)	0.148	0.162	0.141	0.126	0.114	0.102	0.793
(local balance)	0.104	0.195	0.169	0.161	0.172	0.172	0.973
CPU utilization	0.396	0.055	0.050	0.023	0.048	0.046	0.618
	0.493	0.054	0.047	0.021	0.038	0.034	0.687
	0.347	0.065	0.056	0.027	0.057	0.057	0.609
CPU mean queue length	0.40	0.29	0.25	0.21	0.25	0.27	
(simulation)	0.49	0.34	0.34	0.31	0.37	0.38	
CPU standard deviation of queue length	0.49	0.45	0.44	0.41	0.44	0.45	
	0.50	0.47	0.47	0.46	0.48	0.49	
CPU mean wait time	3.33	1.76	1.70	1.53	1.78	1.99	
	3.39	2.06	2.44	2.54	3.16	3.59	
I/O 1 utilization	0.143	0.133	0.060	0.055	0.095	0.092	
	0.169	0.129	0.056	0.047	0.077	0.067	
I/O 2 utilization	0.038	0.106	0.096	0.044	0.076	0.073	
	0.047	0.102	0.087	0.037	0.062	0.062	
I/O 3 utilization	0.063	0.088	0.159	0.220	0.127	0.122	
	0.076	0.093	0.146	0.202	0.101	0.093	
CPU mean service	3.33	0.333	0.333	0.167	0.333	0.333	
I/O 1 probability (mean service = 2.00)	0.6	0.4	0.2	0.2	0.333	0.333	
I/O 2 probability (mean service = 1.60)	0.2	0.4	0.4	0.2	0.333	0.333	
I/O 3 probability (mean service = 2.67)	0.2	0.2	0.4	0.6	0.333	0.333	

were compared for the locally balanced model with and without coalescing. The differences were never more than one percent and usually were less than that.

These programs are more than an order of magnitude faster than existing implementation of other techniques.

8. Conclusions

We have presented approximate solution techniques for several classes of models which are very important in the modeling of computing systems. These techniques are computationally very inexpensive and of great practical value. They complement previous techniques which may be more accurate but are computationally more expensive, and may be used directly in conjunction with previous techniques.

Our techniques give exact results for several interesting classes of models, and are reasonably accurate for typical models of computing systems. Our techniques have been validated extensively. Our methods may be extended to consider more general networks. Our techniques are compatible with the techniques of Keller and Chandy [18] for including the effects of passive re-

sources in central server models. Williams and Bhandiwad [19] have also used approximations of Norton's theorem in analyzing three-class, preemptive, exponential models.

Three sets of programs were used in constructing and validating these models: 1) programs to approximate central server models by two-queue networks, 2) programs to analyze the resulting two-queue networks, and 3) the Crane-Inglehart simulator. Four variations of the two-queue analysis technique were programmed: 1) single class nonexponential, 2) multiple class FCFS, 3) multiple class preemptive priority, and 4) multiple class nonpreemptive priority. Each of these cases (except the last two) required slightly different programs to construct two-queue approximations of the given central server problem. The simulator handles arbitrarily interconnected networks, a large number of customer classes and customers, and a variety of disciplines.

Acknowledgments

The problem was first suggested in discussions with U. Herzog and L. Woo of IBM Research. We are grateful

Table 8 Nonpreemptive, 6 classes, 1 customer/class, 4 I/O devices.

Class	1	2	3	4	5	6	All
CPU throughput	0.211	0.354	0.435	0.384	0.327	0.326	2.036
(simulation)	0.239	0.381	0.408	0.348	0.288	0.270	1.934
(local balance)	0.184	0.334	0.529	0.438	0.368	0.373	2.226
CPU utilization	0.211	0.236	0.036	0.032	0.027	0.022	0.564
	0.239	0.254	0.034	0.029	0.024	0.018	0.598
	0.184	0.223	0.044	0.037	0.031	0.025	0.544
CPU mean queue length	0.23	0.33	0.25	0.19	0.16	0.16	
(simulation)	0.29	0.36	p.27	0.26	0.24	0.25	
CPU standard deviation of queue length	0.42	0.47	0.43	0.39	0.37	0.36	
	0.46	0.48	0.44	0.44	0.43	0.43	
CPU mean wait time	1.08	0.963	0.508	0.478	0.487	0.480	
	1.24	0.946	0.641	0.737	0.857	0.897	
I/O 1 utilization	0.263	0.071	0.087	0.096	0.163	0.163	
	0.305	0.067	0.083	0.081	0.146	0.128	
I/O 2 utilization	0.026	0.106	0.174	0.048	0.082	0.081	
	0.030	0.112	0.177	0.046	0.070	0.074	
I/O 3 utilization	0.013	0.053	0.087	0.024	0.041	0.041	
	0.015	0.057	0.082	0.024	0.038	0.035	
I/O 4 utilization	0.026	0.106	0.043	0.240	0.082	0.081	
	0.028	0.114	0.041	0.216	0.067	0.070	
CPU mean service	1.00	0.667	0.083	0.083	0.083	0.067	
I/O 1 probability (mean service = 2.00)	0.625	0.100	0.100	0.125	0.250	0.250	
I/O 2 probability (mean service = 1.00)	0.125	0.300	0.400	0.125	0.250	0.250	
I/O 3 probability (mean service = 0.50)	0.125	0.300	0.400	0.125	0.250	0.250	
I/O 4 probability (mean service = 1.00)	0.125	0.300	0.100	0.625	0.250	0.250	

for their continued encouragement and criticism. This research was supported in part by National Science Foundation grants GJ-1084 and GJ-35109.

References

1. F. Baskett, "Mathematical Models of Multiprogrammed Computer Systems," TSN-17, Computation Center, The University of Texas, Austin, Texas, January 1971.
2. J. Buzen, "Queueing Network Models of Multiprogramming," Ph.D. Thesis, Division of Engineering and Applied Physics, Harvard University, Cambridge, Massachusetts, June 1971.
3. C.C. Lee, "Queueing Models of Device Utilization in Multiprogrammed Computer Systems," TR-7, Department of Computer Sciences, The University of Texas, Austin, Texas, December 1972.
4. G. S. Shedler, "A Cyclic Queue Model of A Paging Machine," *Research Report RC 2814*, IBM Thomas J. Watson Research Center, Yorktown Heights, New York, March 1970.
5. J. L. Smith, "An Analysis of Time Sharing Computer Systems Using Markov Models," *Proc. Spring Joint Computer Conference (AFIPS)*, 28, Spartan Books, New York, 1966.
6. F. Baskett, K. M. Chandy, R. R. Muntz, and F. Palacios Gomez, "Open, Closed, and Mixed Networks of Queues with Different Classes of Customers," *J. ACM* (to be published).
7. K. M. Chandy, "The Analysis and Solutions for General Queueing Networks," *Proceedings of the Sixth Annual Princeton Conference on Information Sciences*, Princeton University, Princeton, New Jersey, March 1972.
8. D. S. Johnson, "A Process-Oriented Model of Resource Demands in Large, Multiprocessing Computer Utilities," TSN-29, Computation Center, The University of Texas, Austin, Texas, August 1972.
9. K. M. Chandy, U. Herzog, and L. Woo, "Approximate Analysis of General Queueing Networks," *IBM. J. Res Develop.* 19, 43 (1975).
10. V. L. Wallace and Richard S. Rosenberg, "Markovian Models and Numerical Analysis of Computer System Behavior," *Proc. Spring Joint Computer Conference (AFIPS)*, 28, Spartan Books, New York, 1966.
11. M. A. Crane and D. I. Iglehart, "Simulating Stable Stochastic Systems, I: General Multiserver Queues," *J. ACM* 21, 103 (1974).
12. M. A. Crane and D. I. Iglehart, "Simulating Stable Stochastic Systems, II: Markov Chains," *J. ACM* 21, 114 (1974).

13. K. M. Chandy, U. Herzog, and L. Woo, "Parametric Analysis of Queuing Network Models," *IBM J. Res. Develop.* **19**, 36 (1975).
14. U. Herzog, L. Woo, and K. M. Chandy, "Solution of Queuing Problems by a Recursive Technique," *IBM J. Res. Develop.* **19**, 295 (1975), this issue.
15. D. R. Cox, "A Use of Complex Probabilities in the Theory of Stochastic Processes," *Proc. Cambridge Phil. Soc.* **51**, 313 (1955).
16. C. H. Sauer, "Configuration of Computing Systems: An Approach Using Queuing Network Models," Ph.D. Thesis, University of Texas, Austin, Texas, 1975.
17. C. H. Sauer and K. M. Chandy, "Approximate Analysis of Computer System Models with different Classes of Customers," presented at ORSA/TIMS Puerto Rico Meeting, Oct. 1974.
18. T. W. Keller and K. M. Chandy, "Computer Models with Constrained Parallel Processors," presented at Sagamore Conference on Parallel Processing, 1974.
19. A. C. Williams and R. Bhandiwad, Mobil Oil Corp., Princeton, New Jersey, private communication.

Received November 10, 1974

C. H. Sauer is located at the IBM Thomas J. Watson Research Center, Yorktown Heights, New York 10598. K. M. Chandy is located at the Department of Computer Sciences, The University of Texas, Austin, Texas 78712.