

T. W. Gay
P. H. Seaman

Composite Priority Queue

Abstract: This paper presents formulas for calculating waiting time for customers in a queue with combined preemptive and head-of-line (nonpreemptive) priority scheduling disciplines and describes the reasoning behind them. This work has been applied in the development of programmable terminal control units.

Introduction

The development of programmable terminal control units has created the need for a new priority queuing model, one taking into consideration both preemptive and head-of-line priority relationships between devices and programming tasks. As concurrent service demands occur, extended waiting delays and service times are experienced by low priority demands. The effect of higher priority demands on the mean and standard deviation of response time at a particular demand point is presented here.

Previous work on priority queues has been done by Cobham [1], Gaver [2], Takacs [3], Chang [4], Jaiswal [5], and Herzog [6]. Our model is an adaptation of preemptive and head-of-line priority queues, combining the separate results into a composite.

In this paper the model is first characterized, and then the mathematics is developed using heuristic reasoning. Finally, an example is presented applying the model to a typical control unit configuration.

Mathematical description

Suppose that customers of different priorities are arriving at a counter in accordance with a Poisson process of density Λ . The customers are served by a single server in order of priority and for each priority in order of arrival. Each priority consists of two parts:

1. *Preemptive-resume* If a customer of higher preemptive priority arrives when a customer of lower preemptive priority is being served, the server interrupts the current service and immediately starts serving the customer of higher priority. The service of the customer of lower priority is resumed when no more cus-

tomers of higher preemptive priority are present in the system. The server is busy as long as there are customers in the system. There are n levels of such preemptive priority.

2. *Head-of-line (nonpreemptive)* Within each preemptive priority level, m classes of customers have head-of-line priority with respect to each other. If a customer of higher head-of-line priority arrives when a customer of lower priority is being served, the server does not interrupt the current service, as long as both customers have the same preemptive priority level. At the completion of the current service, the server chooses the customer who arrived first among the customers with highest head-of-line priority present in the system.

This composite priority scheme is diagrammed in Fig. 1.

It is convenient to assume that a customer with a smaller priority number has precedence over a customer with a greater priority number. Accordingly, a customer in the j th head-of-line priority class within the i th preemptive priority level has priority number ij . The following notation is used. Let

λ_{ij} = arrival rate of customers with priority number ij
 $\Lambda_i = \sum_{k=1}^i \sum_{j=1}^m \lambda_{kj}$, aggregated arrival rate of customers down to and including the i th preemptive priority level ($\Lambda_n = \Lambda$)

ω_{ij} = mean service time for a customer of priority number ij

$\omega_{ij}^{(r)}$ = r th moment of service time with priority number ij . Here, as well as below, the first moment

is denoted by omission of the superscript. Thus,
 $\omega_{ij}^{(1)} \equiv \omega_{ij}$,
 $a_i^{(r)} = \Lambda_i^{-1} \sum_{k=1}^{i-1} \sum_{j=1}^m \lambda_{kj} \omega_{kj}^{(r)}$, r th moment of aggregated service times down to and including the i th preemptive priority level

$\rho_{ij} = \lambda_{ij} \omega_{ij}$, utilization of the server due to customers of priority number ij

$U_{ij} = \sum_{k=1}^{i-1} \sum_{\ell=1}^m \rho_{k\ell} + \sum_{\ell=1}^j \rho_{i\ell}$, utilization of the server due to all customers down to the i th preemptive priority level, including those in the i th level across to and including the j th head-of-line class.

We wish to determine the total time in the system experienced by a customer of priority number ij . This time consists of a wait for service, followed by the service itself, subject to interruptions. The first and second moments of these quantities are defined as:

W_{ij} = mean wait time for a customer of priority number ij

$W_{ij}^{(2)}$ = second moment of wait time with priority number ij

T_{ij} = mean service time for a customer of priority number ij , extended by interruptions due to higher priority customers

$T_{ij}^{(2)}$ = second moment of extended service time with priority number ij

Q_{ij} = mean time in system for a customer of priority number ij

σ_{ij}^2 = variance of time in system for a customer of priority number ij .

The service time is independent of the preceding wait time. Therefore, the two can be simply combined to determine the total time in the system. The mean and variance resulting from this convolution would be,

$$Q_{ij} = W_{ij} + T_{ij}; \quad (1)$$

$$\sigma_{ij}^2 = (W_{ij}^{(2)} - W_{ij}^2) + (T_{ij}^{(2)} - T_{ij}^2). \quad (2)$$

Heuristic reasoning

It appears that to start from basic probabilities and develop the required distributions for this model would be very laborious. Consequently, it is desirable to use available results for simpler models and modify them to fit this case, if possible. It is important to emphasize that if the reasoning used is valid, the results will be the same as those obtained from the more involved derivation from "first principles."

In ascertaining what simpler models may be applicable, four points are basic to the argument.

1. A customer of priority number ij is not affected by preemptive levels below him (i.e., greater than i), so those levels do not appear in the equations.

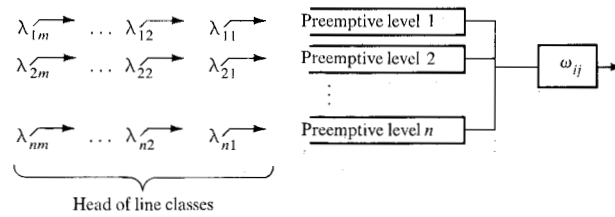


Figure 1 Composite priority scheme.

2. A customer is affected by one set of priority classes during his wait time and another set during service, so these periods must be treated separately. However, because wait and service times remain independent of each other, they may be simply combined, as is done in Eqs. (1) and (2).
3. During the wait time of a customer with priority number ij , all arrivals in preemptive levels above him (i.e., less than i) appear simply as arrivals to an additional head-of-line priority class above him (say class $j=0$). While it is true that this additional priority class is preemptive instead of nonpreemptive, this affects only the order of service of the customers ahead of customer ij , not the length or number of services. As far as customer ij is concerned, his wait is the same regardless of whether the arrivals with priority greater than his are preemptive or nonpreemptive. This was pointed out by Gaver [2].
4. During the service time of a customer of priority number ij , delays are caused by interruptions due to arrivals in preemptive priority levels above him (i.e., less than i). However, arrivals on his own level do not affect him because of the head-of-line discipline within the level. As before, when several interruptions occur concurrently, the order in which they are processed does not concern the customer interrupted because he cannot continue until all have been served one way or another.

Mathematical development

Following Takacs [3], the first two moments of wait time for a customer of priority j in a system including m classes of head-of-line priorities is given by:

$$W_j = \frac{\sum_{k=1}^m \lambda_k \omega_k^{(2)}}{2(1 - U_j)(1 - U_{j-1})}; \quad (3)$$

$$W_j^{(2)} = \frac{\sum_{k=1}^m \lambda_k \omega_k^{(3)}}{3(1 - U_j)(1 - U_{j-1})^2} + \frac{\sum_{k=1}^m \lambda_k \omega_k^{(2)} \sum_{\ell=1}^j \lambda_\ell \omega_\ell^{(2)}}{2(1 - U_j)^2(1 - U_{j-1})^2} + \frac{\sum_{k=1}^m \lambda_k \omega_k^{(2)} \sum_{\ell=1}^{j-1} \lambda_\ell \omega_\ell^{(2)}}{2(1 - U_j)(1 - U_{j-1})^3}, \quad (4)$$

supposing there are no preemptive priority levels above this one and eliding the i subscript.

Table 1 Configuration for a typical control unit.

Interrupt level no.	ij	λ	ω	Erlang- m	ρ	U
1. Cycle stealing for:						
Function a	11	15.4	0.0087	100	0.134	0.134
b	12	0.28	0.014	100	0.004	0.138
c	13	2.1	0.010	100	0.021	0.159
d	14	0.008	0.90	100	0.007	0.166
2. Line control	21	0.35	0.383	100	0.134	0.300
3. Keyboard 1	31	0.005	1.40	2	0.007	0.307
Keyboard 2	32	0.005	1.40	2	0.007	0.314
Keyboard 3	33	0.005	1.40	2	0.007	0.321
Keyboard 4	34	0.005	1.40	2	0.007	0.328
4. Program	41	0.026	4.0	1	0.104	0.432
5. Printer	51	0.022	4.57	2	0.100	0.532

Jaiswall [5] gives this also as Eq. V.6.28; however, as printed the exponent 2 is missing in the denominator of the first term of $W^{(2)}$.

By using point 3) from the reasoning above, the required wait time may be obtained by adding an additional priority class containing all the higher level preemptive traffic from level 1 to $i-1$. This includes adding the quantities $\Lambda_{i-1}a_{i-1}^{(2)}$ or $\Lambda_{i-1}a_{i-1}^{(3)}$ to the various terms in the numerator and $U_{i-1,m}$ to the utilizations in the denominators. Thus,

$$W_{ij} = \frac{\Lambda_i a_i^{(2)}}{2(1 - U_{ij})(1 - U_{i,j-1})}, \quad (5)$$

$$W_{ij}^{(2)} = \frac{\Lambda_i a_i^{(3)}}{3(1 - U_{ij})(1 - U_{i,j-1})^2} + \frac{\Lambda_i a_i^{(2)}[\Lambda_{i-1}a_{i-1}^{(2)} + \sum_{k=1}^j \lambda_{ik} \omega_{ik}^{(2)}]}{2(1 - U_{ij})^2(1 - U_{i,j-1})^2} + \frac{\Lambda_i a_i^{(2)}[\Lambda_{i-1}a_{i-1}^{(2)} + \sum_{k=1}^{j-1} \lambda_{ik} \omega_{ik}^{(2)}]}{2(1 - U_{ij})(1 - U_{i,j-1})^3}. \quad (6)$$

Next the service time for a customer with priority ij , including interruptions, is to be calculated. First, consider a service distribution $H(t)$ without interruptions and denote its transform by

$$\Omega(s) = \int_0^\infty e^{-st} dH(t).$$

The r th moment of service is then given by $\omega^{(r)} = (-1)^r \Omega^{(r)}(0)$. Now let this basic service time be subject to a Poisson stream of interruptions with density λ . Assume that the elapsed time to clear each interruption has a distribution $G(t)$, whose transform is given by

$$A(s) = \int_0^\infty e^{-st} dG(t),$$

with r th moment $a^{(r)} = (-1)^r A^{(r)}(0)$.

The transform of the basic service time, extended by these interruptions, is then expressed as

$$\Omega^*(s) = \Omega\{\lambda[1 - B(s)] + s\} \quad \text{where } B(s) = A\{\lambda[1 - B(s)] + s\}. \quad (7)$$

For the complete argument, see Jaiswal [5], pp 6-12. The random variable represented by $B(s)$ is the busy period for the interruptions.

Let the r th moment of the extended service time be defined as $T^{(r)} = (-1)^r \Omega^{*(r)}(0)$. Then from Eq. (7) we obtain

$$T = \frac{\omega}{1 - \lambda a}, \quad (8)$$

$$T^{(2)} = \frac{\omega^{(2)}}{(1 - \lambda a)^2} + \frac{\omega \lambda a^{(2)}}{(1 - \lambda a)^3}. \quad (9)$$

Applying point 4) from the reasoning above, we may translate these results into the notation of our model as follows:

$$T_{ij} = \frac{\omega_{ij}}{1 - u_{i-1,m}}; \quad (10)$$

$$T_{ij}^{(2)} = \frac{\omega_{ij}^{(2)}}{(1 - u_{i-1,m})^2} + \frac{\omega_{ij} \Lambda_{i-1} a_{i-1}^{(2)}}{(1 - u_{i-1,m})^3}. \quad (11)$$

With this, the development is complete. Should higher moments be required, they may easily be obtained by calculating the additional moments from the basic transforms and applying the modifications as above.

Example Consider the configuration in Table 1 for a typical control unit. Find the responses for the first and last keyboard on interrupt level no. 3. Note that for an Erlang- m distribution with mean T , the second moment is

$\left(1 + \frac{1}{m}\right)T^2$ and the third moment is $\left(1 + \frac{2}{m}\right)\left(1 + \frac{1}{m}\right)T^3$.

$$\Lambda_1 a_1^{(2)} = 0.0012 + 0.00006 + 0.00021 + 0.0063 \\ = 0.0078,$$

$$\Lambda_2 a_2^{(2)} = 0.0078 + 0.052 = 0.060,$$

$$\Lambda_3 a_3^{(2)} = 0.06 + (0.015 + 0.015 + 0.015 + 0.015) \\ = 0.120,$$

$$\Lambda_3 a_3^{(3)} = (0.00001 + 0.000001 + 0.000002 + 0.0057) \\ + 0.0203 + (0.041 + 0.041 + 0.041 + 0.041) \\ = 0.190,$$

$$W_{31} = \frac{0.120}{2(1 - 0.307)(1 - 0.300)} = 0.124,$$

$$W_{31}^{(2)} = \frac{0.190}{3(1 - 0.307)(1 - 0.300)^2} \\ + \frac{0.120(0.06 + 0.015)}{2(1 - 0.307)^2(1 - 0.300)^2} \\ + \frac{0.120(0.06)}{2(1 - 0.307)(1 - 0.300)^3} \\ = 0.187 + 0.019 + 0.015 = 0.221,$$

$$T_{31} = \frac{1.4}{(1 - 0.300)} = 2.0,$$

$$T_{31}^{(2)} = \frac{2.1}{(1 - 0.300)^2} + \frac{1.4(0.06)}{(1 - 0.300)^3} = 4.3 + 0.245 \\ = 4.55.$$

Response time for keyboard 1 is 2.124 ms. Standard deviation of response time is $[(0.221 - 0.015) + (4.55 - 4)]^{\frac{1}{2}} = 0.87$ ms.

$$W_{34} = \frac{0.120}{2(1 - 0.328)(1 - 0.321)} = 0.132$$

$$W_{34}^{(2)} = \frac{0.190}{3(1 - 0.328)(1 - 0.321)^2} \\ + \frac{0.120(0.120)}{2(1 - 0.328)^2(1 - 0.321)^2} \\ + \frac{0.120(0.06 + 0.045)}{2(1 - 0.328)(1 - 0.321)^3} \\ = 0.205 + 0.035 + 0.030 = 0.270$$

$$T_{34} = T_{31}$$

$$T_{34}^{(2)} = T_{31}^{(2)}$$

Response time for keyboard 4 is 2.132 ms.

Standard deviation of response time is $[(0.270 - 0.017) + (4.55 - 4)]^{\frac{1}{2}} = 0.91$ ms. If the 95th percentile of response time is assumed to be at two standard deviations above the mean, then:

Keyboard	95th percentile
1	3.86 ms
4	3.95 ms

Note that the difference here is very slight because the traffic load in interrupt level no. 3 is small. Notice also that the activity in levels 4 and 5 do not affect the calculations at level 3.

Conclusions

Formulas have been presented for a composite priority discipline. An implicit assumption is that interruptions incur no extra overhead (or "service orientation time") when a service is actually preempted. In the areas where the model has been applied, this has been the case, or practically so. However, if service orientation is significant, a more complex model is required. This situation has been studied by Gaver [7] and Jaiswal [5].

References

1. A. Cobham, "Priority Assignment in Waiting Lines," *Oper. Res.* 2, 70 (1954).
2. D. P. Gaver, Jr., "A Waiting Line with Interrupted Service, Including Priorities," *J. Roy. Statist. Soc. B* 24, 73 (1962).
3. L. Takacs, "Priority Queues," *Oper. Res.* 12, 63 (1964).
4. W. Chang, "Queuing with Nonpreemptive and Preemptive-resume Priorities," *Oper. Res.* 13, 1020 (1965).
5. N. K. Jaiswal, *Priority Queues*, Academic Press, New York, 1968.
6. U. Herzog, "Efficient Priority Strategies for Switching Centers in Communication Networks," *Proceedings of the Second Texas Conference on Computing Systems*, University of Texas, Austin, Texas, 1973, p. 45.1.
7. D. P. Gaver, Jr., "Competitive Queuing: Idleness Probabilities Under Priority Disciplines," *J. Roy. Statist. Soc. B* 25, 489 (1963).

Received May 22, 1974

The authors are located at the IBM System Development Division laboratory, Poughkeepsie, New York 12602.