



MATHEMATICAL ANALYSIS OF FEEDBACK QUEUE NETWORK MODEL WITH PRIORITYV COMPRISED OF TWO SERIAL CHANNELS WITHIN STOCHASTIC CONDITIONS

Preeti¹, Deepak Gupta², Vandana Saini³

^{1,2}Department of Mathematics and Humanities, MMEC-Maharishi
Markandeshwar (Deemed to be University) Mullana (Ambala), 133207, India.

³Department of Mathematics, Govt. P.G. College Naraingarh (Ambala),
134203, India.

Email: ¹raziagolu@gmail.com, ²guptadeepak2003@gmail.com
³sainivandana003@gmail.com

Corresponding Author: **Preeti**

<https://doi.org/10.26782/jmcms.2025.07.00009>

(Received: April 13, 2025; Revised: June 26, 2025; Accepted: July 08, 2025)

Abstract

This paper presents a comprehensive analysis of a feedback queue model with a Priority mechanism and investigates its behavior under stochastic conditions. This model comprises two serially connected service channels, with priority applied exclusively to the first service channel. Upon entry, customers are classified into two groups-low and high priority. A preemptive priority discipline is used at the first server to distinguish between high- and low-priority customers, thereby reflecting real-world service hierarchies. The feedback mechanism in the model allows for a maximum of one time only for the customer's satisfaction with the service. The arrival of the customers is governed by a Poisson process and and service times at both servers are assumed to follow independently and be exponentially distributed. Upon service completion at the second server, customers may either exit the system permanently or re-enter the network through a feedback loop. The Steady-state behavior of the system is captured through a set of differential equations, which are solved by using the generating function technique combined with classical calculus laws. Various queue performance indicators, including average queue length, variance in queues, server utilization, and total duration time, are discussed. In the last section, a comparative study of the model with the literature is also discussed. The model's behaviour is well demonstrated both graphically and numerically and provides an in-depth understanding of how each parameter influences the overall system performance, and the obtained results prove the stability and accuracy of the model. The insights derived from the analysis could help understand the design and optimization of the queueing model in different settings such as hospitals, manufacturing industries, and telecommunications.

Preeti et al.

Keywords: Feedback, Generating function techniques, Priority, Queueing, Serial Channel, Stochastic condition.

I. Introduction

Queueing theory plays a vital role in the analysis of diverse domains such as industrial systems, manufacturing companies, hospitals, and telecommunications networks. In the queueing theory context, it is necessary to identify that not all users are always homogeneous; they can exhibit heterogeneous characteristics. Due to these differences, it is essential to implement a priority mechanism in the system. O'Brien [XIII] discussed the problems of queueing theory with their potential Solution. Jackson [VII] contributed to studying the queueing problems for two-phase and three-phase systems. Finch P.D. [IV] expanded the field further by examining the cyclic queues with feedback. Subsequent advancements in queueing theory have included the contributions of Luo C. et al [XII], who discussed the transient queue size distribution solution of queues with feedback, and Pang Y. [XIV] introduced the concept of a retrial queue system with preemptive resume and Bernoulli Feedback. Later on, Singh T.P et al [XIX] made a remarkable work in analyzing the behavior of various feedback queue models, and Kusum [VIII] developed mathematical modeling of various feedback and fuzzy queue network models. Tyagi A. et al. [XX] explained the behavior of impatient customers by taking a serial queue model in a stochastic environment. Zadeh A.B. [XXII] focused on a multi-phase system including random feedback in service, single vacation, and batch arrival. Kumar S.. et al. [IX- X] studied a feedback queueing model of three servers and explored the probability of customers receiving service at most twice. Gupta R. et al [V-VI] analyzed a bi-serial server with bulk arrival, contributing to the understanding of complex queueing networks. After that, Ajewole O.R et al [I] explored a Preemptive priority queue model in which the service unit follows an Erlang type distribution. Later on, Dudin A. et al [III, XI] further advanced the field by analyzing both single server and multi-class queues with batch arrivals, unreliable service, dynamic change of Priorities, and also analyzed a priority queueing system with enhanced fairness of servers. After that, Saini V. et al [XVI] analyzed the behavior of a feedback queue system consisting of two serial servers, with arriving customers are impatient in behavior. Saini A. et al [XV] discussed the mathematical study of two serial servers with the concept of priority and reneging. Later on, Sangeeta et al [XVII] analyzed a serial queue with the use of reneging, feedback, and discouragement, offering a more nuanced approach to understanding customer behavior. Agarwal D. et al [II] participated by examining a retrial priority G-queue network model with Bernoulli feedback and detecting the optimal working vacation service rate. Shree V. et al [XVIII] studied a priority queue system with batch arrival and Bernoulli feedback and determined the optimal working vacation service rate. Recently, Liu L. et al [XXI] examined a two-stage Queueing system with priority.

In the current study, we undertake an in-depth analysis of a priority-based feedback queueing model, which extends the foundational work of Saini A. et al. [XV]. The main objective of this paper is to investigate the key queue characteristics using a combination of advanced mathematical tools, including probability-generating functions and calculus-based methodologies. Additionally, graphical analysis is to be

Preeti et al.

performed to facilitate a deeper understanding of the model and to illustrate its behavior easily.

II. Model Description, Notations used, and their Practical implication

Assumptions

The following assumptions of this model are:

- A customer who is willing to get service may join the first service channel, and after getting service from this service channel, they will go to the next phase for service.
- Arrival unit follows a Poisson distribution, and the Service unit follows an exponential distribution.
- Feedback is allowed only once, i.e., Revisit of the customers is allowed only once.
- A Preemptive priority discipline is followed by customers
- The Arrival Population is infinite in this model.

Practical Implications

Queueing theory has diverse and widespread applications across various sectors, offering valuable insights for customer satisfaction and operational efficiency. Queueing theory has been used in banking, telecommunications, hospitals, parks, offices, manufacturing industries, educational institutions, saloons, etc. One practical example can be seen in the context of a theater. In this setting, customers seeking to purchase tickets can be divided into two groups: low-priority customers like adults, and high-priority customers like kids, VIPs, senior citizens, etc. These groups are processed according to a priority rule; high-priority customers are given precedence in getting service at the ticket counter. In the case where a customer is unsuccessful in getting a ticket during their initial attempts, the system allows revisiting the facility at the ticket counter. The revisit facility not only enhances customer satisfaction but also emphasizes the importance of prioritizing particular customer groups to manage resources effectively. By applying appropriate priority and feedback mechanisms, service providers can better balance customer demand with available resources, reduce wait duration, and increase overall satisfaction.

Notation Used in the Model

The terminology applied in the current study is shown in Table 1:

Table 1(Notations used in the model)

Service Channel	SC ₁	SC ₂
Arrival Rate	$\lambda_{1H}, \lambda_{1L}$	-
Service Rate	μ_{1H}, μ_{1L}	μ_2
No. of Customers	n_{1H}, n_{1L}	n_2
Probability of the Customers moving from one service channel to another service channel	$SC_1 \rightarrow SC_2 \rightarrow a_{12}$	$SC_2 \rightarrow SC_1 \rightarrow a_{21}$ $SC_2 \rightarrow \text{exit} (a_2)$

Preeti et al.

III. Mathematical Modeling of the Proposed Model

In this model, there are two service channels, SC_1 and SC_2 , which are in series. Both types of customers (low and high priority) arrive in the system to get service with arrival rates $\lambda_{1H}, \lambda_{1L}$ and service rates μ_{1H}, μ_{1L} . Priority is allowed only service channel SC_1 . The arriving customers first go to service channel SC_1 to take service based on some priority rule, and after getting service from service channel SC_1 , they may go to service channel SC_2 with a probability a_{12} . After getting service from SC_2 , either a customer exits from the system with probability a_2 , or if he is not satisfied with the service, then he may visit again to the service channel SC_1 again with probability a_{21} . Finally, he may exit from the system with the same probability a_2 demonstrated in Figure 1:

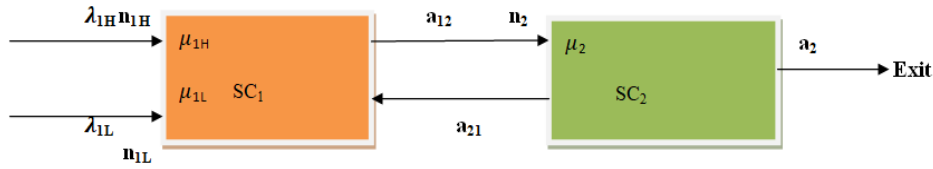


Fig. 1. (Proposed Priority and Feedback Queue Model with two Serial Channels)

Justification for Parameter Selection

The arrival and service rate parameters (such as $\lambda_{1H}, \lambda_{1L}, \mu_{1H}, \mu_{1L}$ and μ_2) utilized in mathematical computations are selected in accordance with typical ranges noticed in similar queueing models [XV, XVI]. These values correspond to moderate traffic intensity ($\lambda/\mu < 1$), thereby ensuring the system stability (i.e system remains in steady state). Their ranges were further evaluated using a sensitivity analysis to ensure the robustness of the performance measures.

IV. Steady State Analysis of the Model

The Steady state form of the differential difference equation is outlined as below:

When $n_{1H}, n_{1L}, n_2 \succ 0$

$$(\lambda_{1H} + \lambda_{1L} + \mu_{1H} + \mu_2)P_{n_{1H}, n_{1L}, n_2} = \lambda_{1H}P_{n_{1H}-1, n_{1L}, n_2} + \lambda_{1L}P_{n_{1H}, n_{1L}-1, n_2} + \mu_{1H}P_{n_{1H}+1, n_{1L}, n_2-1} + \mu_2 a_{21}P_{n_{1H}-1, n_{1L}, n_2+1} + \mu_2 a_2 P_{n_{1H}, n_{1L}, n_2+1} \quad (I)$$

When $n_{1H}, n_{1L}, n_2 = 0$

$$(\lambda_{1H} + \lambda_{1L})P_{0,0,0} = \mu_2 a_2 P_{0,0,1} \quad (II)$$

Solution Methodology

By taking all the possible combinations of n_{1H}, n_{1L}, n_2 , a Total of (I) to (VIII) equations have been obtained. Both Generating and partial generating function techniques are employed to solve the differential difference equation in steady state as given in the following sections:

$$N(X, Y, Z) = \sum_{n_{1H}=0}^{\infty} \sum_{n_{1L}=0}^{\infty} \sum_{n_2=0}^{\infty} P_{n_{1H}, n_{1L}, n_2} X^{n_{1H}} Y^{n_{1L}} Z^{n_2} \quad (1)$$

And the Partial generating function techniques are given as:

$$N_{n_{1L}, n_2}(X) = \sum_{n_{1H}=0}^{\infty} P_{n_{1H}, n_{1L}, n_2} X^{n_{1H}} \quad (2)$$

$$N_{n_2}(X, Y) = \sum_{n_{1L}=0}^{\infty} N_{n_{1L}, n_2}(X) Y^{n_{1L}} \quad (3)$$

$$N(X, Y, Z) = \sum_{n_{1L}=0}^{\infty} N_{n_2}(X, Y) Z^{n_2} \quad (4)$$

The solution for steady-state equations may be obtained by solving equations (I) to (VIII) by using (1) to (4) gives equation (5):

$$N(X, Y, Z) = \frac{\begin{bmatrix} S_1 \left[\mu_{1H} \left(1 - \frac{Z}{X}\right) - \mu_{1L} \left(1 - \frac{Z}{Y}\right) \right] + S_2 \left[\mu_2 \left(1 - \frac{a_{21}X}{Z} - \frac{a_2}{Z}\right) \right] \\ + S_3 \left[\mu_{1L} \left(1 - \frac{Z}{Y}\right) \right] \end{bmatrix}}{\lambda_{1H}(1-X) + \lambda_{1L}(1-Y) + \mu_{1H} \left(1 - \frac{Z}{X}\right) + \mu_2 \left(1 - \frac{a_{21}X}{Z} - \frac{a_2}{Z}\right)} \quad (5)$$

For convenience, let us take

$$S_1 = N_0(Y, Z), S_2 = N_{0,0}(Z), S_3 = N_0(X, Y)$$

Now, after putting $X=Y=Z=1$ in equation (5), an indeterminate form (i.e. $\left[\frac{0}{0} \right]$) arises

whose solution can be found by using L'Hospital's rule. Taking $Y=Z=1$ and differentiating w.r.t X , we get:

$$S_1 \mu_{1H} - a_{21} S_2 \mu_2 = -\lambda_{1H} + \mu_{1H} - a_{21} \mu_2 \quad (6)$$

Taking $X=Z=1$ and differentiating w.r.t Y , we get:

$$-S_1 \mu_{1L} + S_3 \mu_{1L} = -\lambda_{1L} \quad (7)$$

Taking $X=Y=1$ and differentiating w.r.t Z , we get:

$$S_1(-\mu_{1H} + \mu_{1L}) + S_2\mu_2 - S_3\mu_{1L} = -\mu_{1H} + \mu_2 \quad (8)$$

Solving equations (6) to (8), we obtain the following values of S_1 , S_2 , and S_3 :

$$S_1 = 1 - \left(\frac{a_{21}\lambda_{1L} + \lambda_{1H}}{\mu_{1H}(1 - a_{21})} \right) \quad (9)$$

$$S_2 = 1 - \left(\frac{\lambda_{1L} + \lambda_{1H}}{\mu_2(1 - a_{21})} \right) \quad (10)$$

$$S_3 = 1 - \left(\frac{a_{21}\lambda_{1L} + \lambda_{1H}}{\mu_{1H}(1 - a_{21})} + \frac{\lambda_{1L}}{\mu_{1L}} \right) \quad (11)$$

And the Utilization factor at different Servers is given below:

$$\begin{aligned} \nu_1 &= \left(\frac{a_{21}\lambda_{1L} + \lambda_{1H}}{\mu_{1H}(1 - a_{21})} \right) \\ \nu_2 &= \left(\frac{\lambda_{1L} + \lambda_{1H}}{\mu_2(1 - a_{21})} \right) \\ \nu_3 &= \left(\frac{a_{21}\lambda_{1L} + \lambda_{1H}}{\mu_{1H}(1 - a_{21})} + \frac{\lambda_{1L}}{\mu_{1L}} \right) \end{aligned}$$

Convergence and Numerical Stability of Generating Function Approach

In the present Feedback-priority queue network model, the system behavior has been investigated through a sequence of steady-state differential equations (see Equations (I) to (VIII)). These equations were transformed into generating function form using a 3-variable probability generating function (PGF), as defined in Equation (1):

$$N(X, Y, Z) = \sum_{n_{1H}=0}^{\infty} \sum_{n_{1L}=0}^{\infty} \sum_{n_2=0}^{\infty} P_{n_{1H}, n_{1L}, n_2} X^{n_{1H}} Y^{n_{1L}} Z^{n_2}$$

To ensure the numerical validity of this function and the model results derived from it, the convergence of the generating function was checked under the following scenarios:

- **Convergence Region:** For the PGF to be convergent, the magnitude of the variables must fulfill the condition $|X|=|Y|=|Z|\leq 1$
- **Stability Condition:** The system's traffic intensity < 1 for each queueing service channel, guaranteeing that the PGF series converges absolutely.
- **Normalization Check:** At $X=Y=Z=1$, the generating function meets $N(1,1,1)=1$, confirming the total probability condition.
- **Indeterminate form:** At $X=Y=Z=1$ and $N(1,1,1)=1$, an indeterminate form must be required to provide the result.

- Similar convergence methods have been used in related works (Kusum [VIII], Gupta R. [V]).

Queue Performance Indicators

The Queue performance indicators have been calculated by using the following formulas:

a) Expected Queue length:

$$L = L_{n_{1H}} + L_{n_{1L}} + L_{n_2}$$

Where

$$L_{n_{1H}} = \frac{v_1}{1-v_1}, L_{n_{1L}} = \frac{v_3}{1-v_3}, L_{n_2} = \frac{v_2}{1-v_2}$$

b) Variance in Queues:

$$Variance = V_{n_{1H}} + V_{n_{1L}} + V_{n_2}$$

Where

$$V_{n_{1H}} = \frac{v_1}{(1-v_1)^2}, V_{n_{1L}} = \frac{v_3}{(1-v_3)^2}, V_{n_2} = \frac{v_2}{(1-v_2)^2}$$

(c) Expected Wait Time of the customers:

$$Expt(W) = \frac{L}{\lambda}, \text{Where } \lambda = \lambda_{1H} + \lambda_{1L}$$

Behavioral Analysis of the Model through Numerical Illustrations

In the following analysis, different queue performance metrics, including average queue length, wait time, and variance in queues, etc., are examined by taking different combinations of arrival and service rate as given in Table 2:

Table 2: Arrival rate for low and high priority customers, service rate for low and high priority customers, and probability values for Customers' visits

$$\lambda_{1H} = 5, \lambda_{1L} = 3, \mu_{1H} = 19, \mu_{1L} = 21, \mu_2 = 16, a_{21} = 0.2, a_2 = 0.8$$

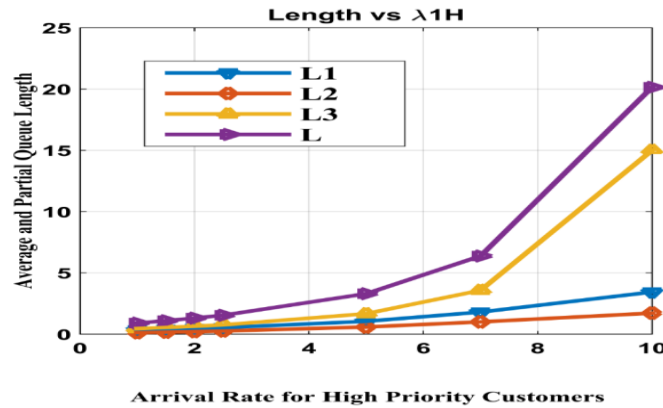


Fig. 1. Average and partial queue length versus arrival rate for high priority customers, i.e., λ_{1H} w.r.t. Table 3

Table 3: Traffic Intensity, Average Queue Length, Partial queue length, customer's waiting time, and variance of queues w.r.t λ_{1H} i.e. arrival rate for high priority customers as shown in the table:

λ_{1H}	ν_1	ν_2	ν_3	L_1	L_2	L_3	L	Variance	Expt (W)
1	0.1053	0.3125	0.2481	0.3300	0.1176	0.4545	0.9022	1.2315	0.1128
1.5	0.1382	0.3516	0.2810	0.3908	0.1603	0.5422	1.0933	1.5657	0.1367
2	0.1711	0.3906	0.3139	0.4575	0.2063	0.6410	1.3049	1.9677	0.1631
2.5	0.2039	0.4297	0.3468	0.5309	0.2562	0.7534	1.5406	2.4557	0.1926
5	0.3684	0.6250	0.5113	1.0462	0.5833	1.6667	3.2962	7.5086	0.4120
7	0.5000	0.7813	0.6429	1.8000	1.0000	3.5714	6.3714	23.3665	0.7964
10	0.6316	0.9375	0.7744	3.4333	1.7143	15	20.1476	259.8742	2.5185

Table 4: Traffic Intensity, Average Queue Length, Partial queue length, customer's waiting time, and variance of queues w.r.t λ_{1L} , i.e., arrival rate for low priority customers as shown in the table:

λ_{1L}	ν_1	ν_2	ν_3	L_1	L_2	L_3	L	Variance	Expt (W)
1	0.3421	0.4688	0.3897	0.6386	0.5200	0.8824	2.0410	3.4977	0.2551
1.5	0.3487	0.5078	0.4201	0.7245	0.5354	1.0317	2.2916	4.1675	0.2864
2	0.3553	0.5469	0.4505	0.8198	0.5510	1.2069	2.5778	5.0101	0.3222
2.5	0.3618	0.5859	0.4809	0.9264	0.5670	1.4151	2.9085	6.0906	0.3636
3	0.3684	0.6250	0.5113	1.0462	0.5833	1.6667	3.2962	7.5086	0.4120
4.5	0.3882	0.7422	0.6024	1.5154	0.6344	2.8788	5.0286	16.0148	0.6286
7	0.4211	0.9375	0.7544	3.0714	0.7273	15	18.7987	253.7613	2.3498

Table 5: Traffic Intensity, Average Queue Length, Partial queue length, customer's waiting time, and variance of queues w.r.t μ_{1H} i.e., service rate for high priority customers as shown in the table:

μ_{1H}	ν_1	ν_2	ν_3	L_1	L_2	L_3	L	Variance	Expt (W)
19	0.3684	0.6250	0.5113	1.0462	0.5833	1.6667	3.2962	7.5086	0.4120
20	0.3500	0.6250	0.4929	0.9718	0.5385	1.6667	3.1770	7.1891	0.3971
21	0.3333	0.6250	0.4762	0.9091	0.5000	1.6667	3.0758	6.9300	0.3845
22	0.3182	0.6250	0.4610	0.8554	0.4667	1.6667	2.9888	6.7161	0.3736
26	0.2692	0.6250	0.4121	0.7009	0.3684	1.6667	2.7360	6.1408	0.3420
30	0.2333	0.6250	0.3762	0.6031	0.3043	1.6667	2.5741	5.8081	0.3218
40	0.1750	0.6250	0.3179	0.4660	0.2121	1.6667	2.3448	5.3847	0.2931

Table 6: Traffic Intensity, Average Queue Length, Partial queue length, customer's waiting time, and variance of queues w.r.t μ_{1L} i.e., service rate for low priority customers as shown in the table:

μ_{1L}	V_1	V_2	V_3	L_1	L_2	L_3	L	Variance	Expt(W)
21	0.3684	0.6250	0.5113	1.0462	0.5833	1.6667	3.2962	7.5086	0.4120
23	0.3684	0.6250	0.4989	0.9954	0.5833	1.6667	3.2454	7.3544	0.4057
26	0.3684	0.6250	0.4838	0.9373	0.5833	1.6667	3.1873	7.1838	0.3984
29	0.3684	0.6250	0.4719	0.8935	0.5833	1.6667	3.1435	7.0598	0.3929
35	0.3684	0.6250	0.4541	0.8320	0.5833	1.6667	3.0820	6.8922	0.3852
43	0.3684	0.6250	0.4382	0.7800	0.5833	1.6667	3.0300	6.7563	0.3787
50	0.3684	0.6250	0.4284	0.7495	0.5833	1.6667	2.9995	6.6794	0.3749

Table 7: Traffic Intensity, Average Queue Length, Partial queue length, customer's waiting time, and variance of queues w.r.t μ_2 i.e., service rate for the second service channel as shown in the table:

μ_2	V_1	V_2	V_3	L_1	L_2	L_3	L	Variance	Expt(W)
16	0.3684	0.6250	0.4284	0.7495	0.5833	1.6667	2.9995	6.6794	0.3749
17	0.3684	0.5882	0.4284	0.7495	0.5833	1.4286	2.7614	5.7043	0.3452
18	0.3684	0.5556	0.4284	0.7495	0.5833	1.2500	2.5829	5.0475	0.3229
30	0.3684	0.3333	0.4284	0.7495	0.5833	0.5000	1.8329	2.9850	0.2291
32	0.3684	0.3125	0.4284	0.7495	0.5833	0.4545	1.7874	2.8961	0.2234
34	0.3684	0.2941	0.4284	0.7495	0.5833	0.4167	1.7495	2.8252	0.2187
36	0.3684	0.2778	0.4284	0.7495	0.5833	0.3846	1.7175	2.7675	0.2147

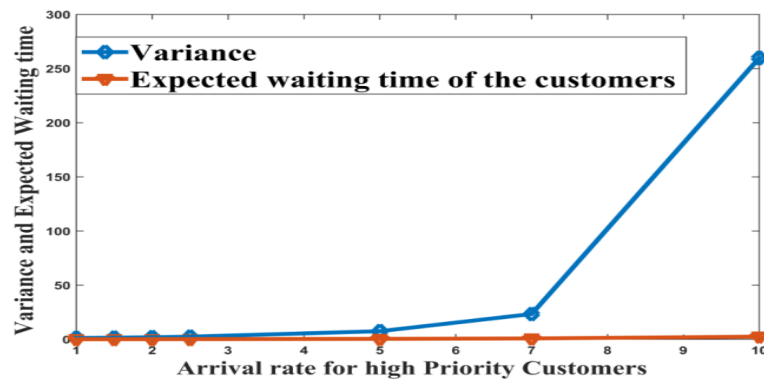


Fig. 3: Expected waiting time and variance of queues versus arrival rate for high priority customers, i.e. λ_{1H} w.r.t Table 3

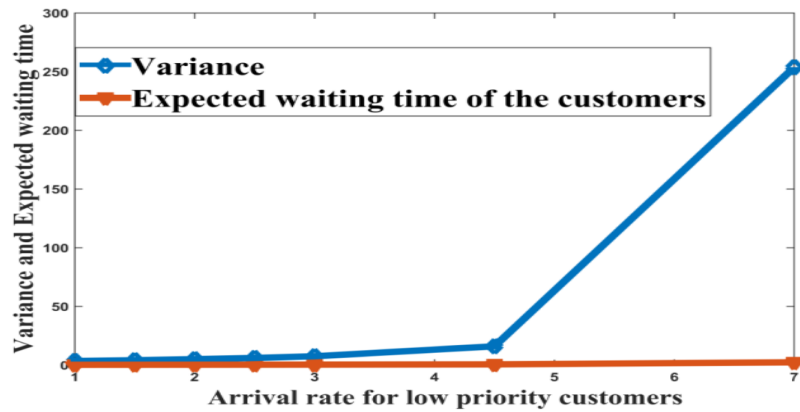


Fig. 4: Expected waiting time and variance of queues versus arrival rate for low priority customers, i.e. λ_{1L} w.r.t Table 4

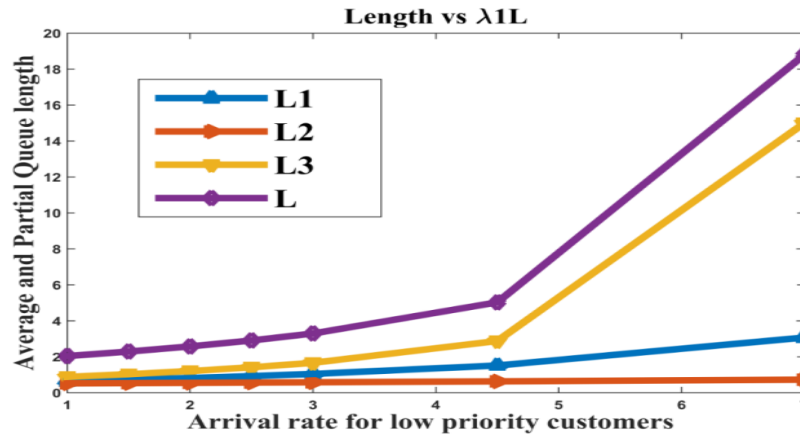


Fig. 5: Average and partial Queue Length versus Arrival rate for low priority customer, i.e. λ_{1L} w. r. t Table 4

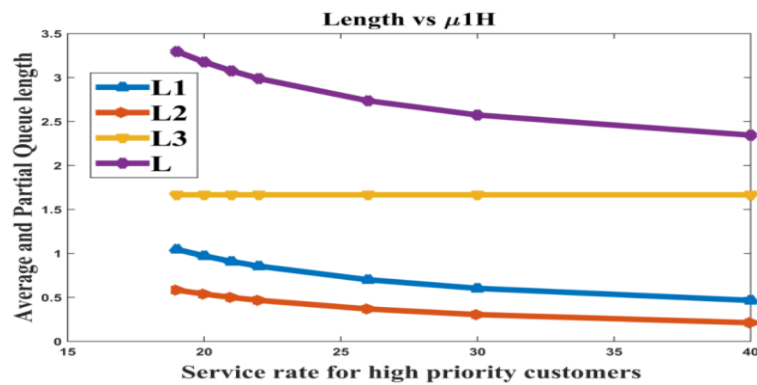


Fig. 6: Average and partial Queue Length versus Service rate for high priority customer, i.e. μ_{1H} w. r. t Table 5

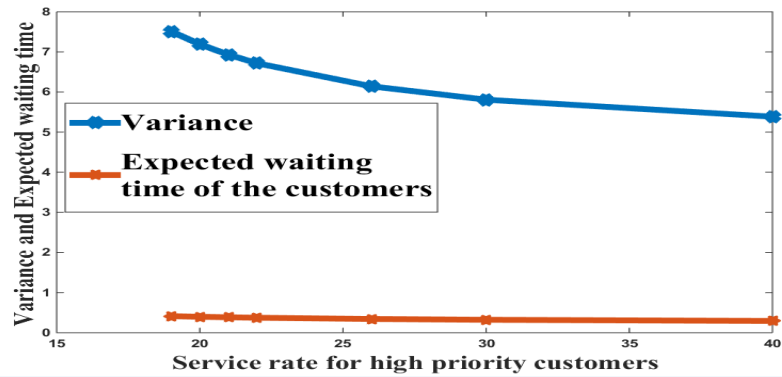


Fig. 7. Expected waiting time and variance of queues versus service rate for high priority customer, i.e. μ_{1H} w.r. t. Table 5

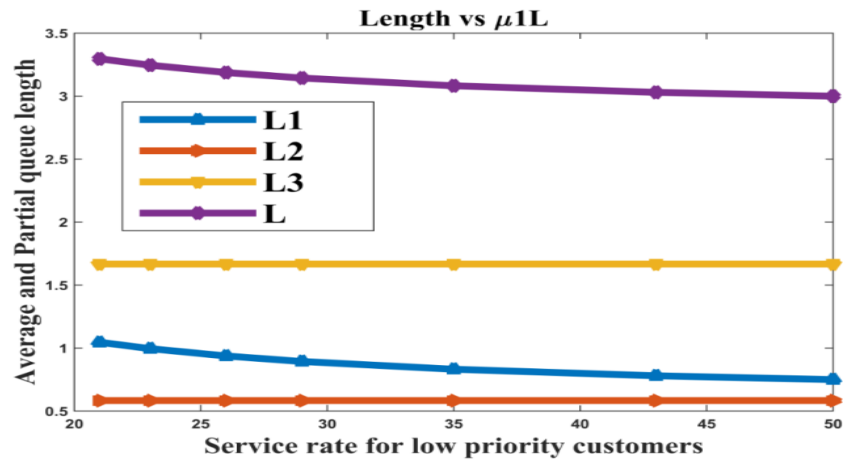


Fig. 8. Average and partial Queue Length versus Service rate for low priority customer, i.e. μ_{1L} w.r.t Table 6

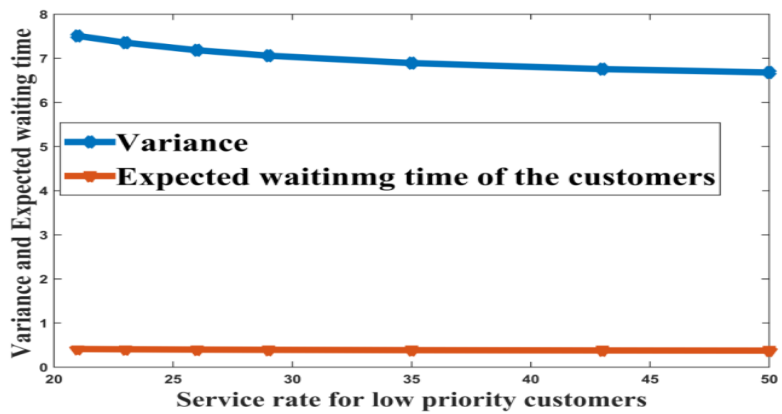


Fig. 9. Expected waiting time and variance of queues versus service rate for low priority customer, i.e. μ_{1L} w.r.t. Table 6

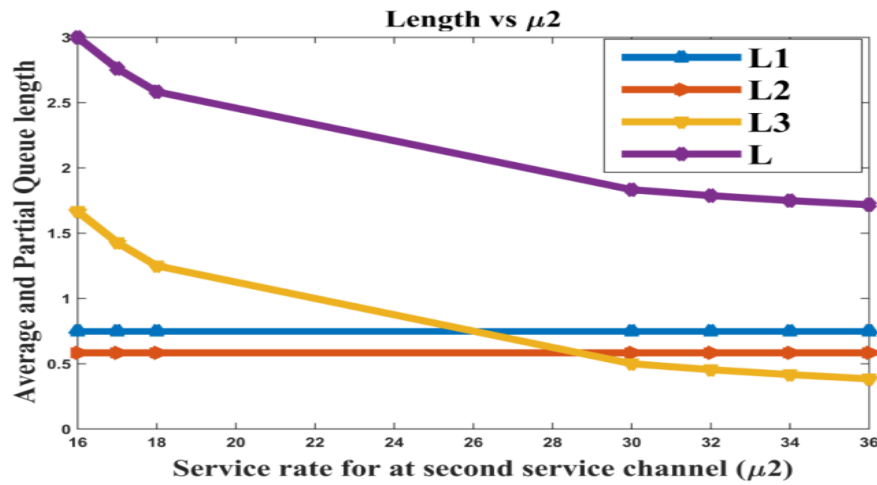


Fig. 10. Average and partial queue length versus Service rate μ_2 w.r.t Table 7

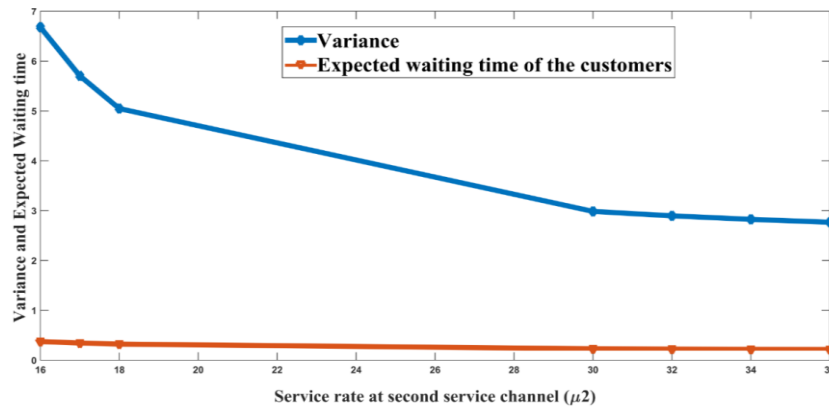


Fig. 11. Expected waiting time and variance of queues versus service rate for the second service channel, i.e. μ_2 W.r.t. Table 7

V. Comparative Study through existing Literature Data

Let us compare the average queue length derived through this model with existing literature data, particularly the work done by Saini A. et al [XV], by taking different service rates for low and high priority customers as used in her study. The following key questions will arise during the comparison of data:

- Does any significant difference exist in average queue length between the model and the literature data?
- What happens if feedback is applied at each of the service channels, particularly when integrated with a priority system?
- How does the service efficiency change when feedback is incorporated alongside a revisiting facility, at most once?

Table 8: Comparison of the Average Queue Length for various values of Service rate, juxtaposing the Literature data with Model data (Calculating the relative differences between Model and Literature Data)

Service rate μ_{1H}	Average queue length(from Model)	Average Queue length(literature data)	Difference b/w the model and the literature data	Service rate μ_{1L}	Average queue length(from Model)	Average Queue length(literature data)	Difference b/w the model and the literature data
10	inf	2.9060	Inf	7	15.2500	2.659	12.591
11	17.7143	2.6559	15.0584	8	9.2500	2.4478	6.8022
12	10.2214	2.4646	7.7568	9	7.2500	2.3003	4.9497
13	7.6071	2.3138	5.2933	10	6.2500	2.1916	4.0584
14	6.2500	2.1914	4.0586	11	5.6500	2.1082	3.5418
15	5.4107	2.0896	3.3211	12	5.2500	2.0424	3.2076
16	4.8373	2.0055	2.8318	13	4.9643	1.9887	2.9756
17	4.4194	1.9330	2.4864	14	4.7500	1.9446	2.8054
18	4.1006	1.8708	2.2298	15	4.5833	1.9078	2.6755

Graphical Representation of the Model and Literature Data w.r.t various service rates for both high and low priority customers:

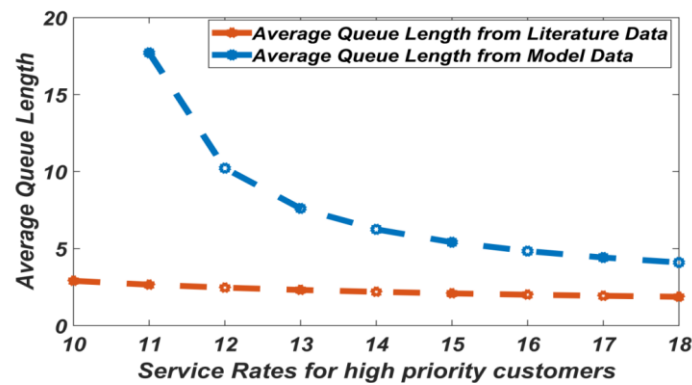


Fig. 12. Comparison of Average Queue Length of the Model and Literature Data by taking various values of service rates for high priority customers

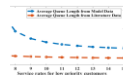


Fig. 13. Comparison of Average Queue Length of Model and Literature Data by taking various values of service rates for low priority customers

VI. Results and Discussion

As we observed that with an increase in the arrival rate for low and high priority customers leads to increase in the partial and average queue length of the system, variance in queues and expected waiting time of the customers, this trend is illustrates in figures 2 to 5 (refer to Tables 3 and 4). Specifically, while slight increase in λ_{IH} and λ_{IL} ranging from 1 to 2.5, and then there is not much increase in these performance metrics. However, when the parameters $\lambda_{IH} = 7$ and $\lambda_{IL} = 4.5$ are applied, a significant increase in all queueing features is observed, indicating a great impact on system performance.

It is evident from Figures 6 to 9 and Tables 5 and 6 that an increase in service rate leads to a reduction in both the partial and average queue length as well as waiting time and variance, but the queue length of the second and third servers remains constant. This effect is most prominent in the first and second servers, where increased service efficiency results in shorter queues. However, the queue length at the third server remains largely unaffected, even as the system as a whole experiences improved performance due to reduced congestion and variability.

Furthermore, as shown in Table 7 and Figures 10 to 11, an increase in the service rate of the system reduces the queue length at the third server, while the queue length at the first and second servers remains constant. The resulting variation leads to a decrease in the average queue length, waiting time, and variance of queues through the entire system. When service rates at the second server increase from 18 to 30, a marked reduction in all queueing features is observed, reflecting improved service efficiency and reducing congestion.

From comparative evaluation with existing models from the literature (refer to Table 8), it is demonstrated that when feedback is applied at all service stages with one one-time revisiting facility for customers, then the mean queue length rises faster as compared to the model that implements priority at the first service channel. As the service rates increase, the average queue length correspondingly decreases, because more and more customers are served within the system. This indicates that the

application of feedback in conjunction with the priority based model improves the efficiency of the system by reducing delays and improving service outputs. The graphical representation of these results is shown in Figures 12 and 13, which describe the effect of feedback on system performance.

VII. Conclusion and Future Scope

The present research investigates a feedback queue model with priority at the first server, and a single revisiting option demonstrates improved system stability over traditional models and reduced average queue length, particularly under moderate traffic intensity. Results highlight performance trade-offs between priority handling and customer feedback mechanisms. Compared to priority-only systems, feedback incorporation improves throughput but may increase queue length under certain conditions. These findings are significant for service environments such as healthcare triage systems and multi-stage call centers, where prioritization and limited re-entry are operationally common. Comparative analysis of the model with literature data has been conducted in the last section to prove the model's validation and accuracy. Future work may investigate biserial queue configurations, multiple feedback loops, or fuzzy environments to further generalize the applicability of this approach in real-world systems such as telecommunications and healthcare networks. The model can be extended to include multiple service channels in parallel, in addition to serial ones, to better reflect real-world systems.

Conflicts of Interest:

There were no relevant conflicts of interest in this Paper.

References

- I. Ajewole, O. R., C. O. Mmduakor, E. O. Adeyefa, J. O. Okoro, and T. O. Ogunlade. "Preemptive-Resume Priority Queue System with Erlang Service Distribution." *Journal of Theoretical and Applied Information Technology*, vol. 99, no. 6, 2021, pp. 1426–1434.
- II. Agarwal, A., R. Agarwal, and S. Upadhyaya. "Detection of Optimal Working Vacation Service Rate for Retrial Priority G-Queue with Immediate Bernoulli Feedback." *Results in Control and Optimization*, vol. 14, 2024, p. 100397. 10.1016/j.rico.2024.100397.
- III. Dudin, A., O. Dudina, S. Dudin, and K. Samouylov. "Analysis of Single-Server Multi-Class Queue with Unreliable Service, Batch Correlated Arrivals, Customer's Impatience and Dynamical Change of Priorities." *Mathematics*, vol. 9, no. 1257, 2021, pp. 1–17. 10.3390/math9111257.
- IV. Finch, P. D. "Cyclic Queues with Feedback." *Journal of the Royal Statistical Society: Series B (Methodological)*, vol. 21, no. 1, 1959, pp. 153–157. <https://www.jstor.org/stable/i349709>.

- V. Gupta, Renu. Queue Network Models and Their Applications. 2023. Maharishi Markandeshwar University, Mullana, PhD dissertation. <http://hdl.handle.net/10603/472701>.
- VI. Gupta, R., and D. Gupta. "Steady State Analysis of Biseria Server with Bulk Arrival." *International Journal of Mathematics Trends and Technology*, vol. 56, no. 6, 2018, pp. 430–436.
- VII. Jackson, R. R. P. "Queuing System with Phase Type Services." *O.R. Quarterly*, vol. 5, 1954, pp. 109–120. 10.2307/3007088.
- VIII. Kusum. Mathematical Modeling of Feedback and Fuzzy Queue Network. PhD thesis, Maharishi Markandeshwar University, Department of Mathematics, 2012. <http://hdl.handle.net/10603/10532>.
- IX. Kumar, S., and G. Taneja. "A Feedback Queuing Model with Chances of Providing Service at Most Twice by One or More out of Three Servers." *Aryabhatta Journal of Mathematics and Informatics*, vol. 9, no. 1, 2017, pp. 171–184.
- X. Kumar, S., and G. Taneja. "A Feedback Queueing Model with Chance of Providing Service at Most Twice by One or More out of Three Servers." *International Journal of Applied Engineering Research*, vol. 13, no. 17, 2018, pp. 13093–13102.
- XI. Lee, S., A. Dudin, and O. Dudina. "Analysis of a Priority Queueing System with the Enhanced Fairness of Servers Scheduling." *Journal of Ambient Intelligence and Humanized Computing*, vol. 15, no. 1, 2024, pp. 465–477. <https://doi.org/10.1007/s12652-022-03903-z>.
- XII. Luo, C., Y. Tang, and C. Li. "Transient Queue Size Distribution Solution of Geom/G/1 Queue with Feedback, a Recursive Method." *Journal of Systems Science and Complexity*, vol. 22, 2009, pp. 303–312. 10.1007/s11424-009-9165-7.
- XIII. O'Brien, G. G. "The Solution of Some Queueing Problems." *Journal of the Society for Industrial and Applied Mathematics*, vol. 2, no. 3, 1954, pp. 133–142. 10.1137/0102010.
- XIV. Peng, Y. "On the Discrete-Time Geo/G/1 Retrial Queueing System with Preemptive Resume and Bernoulli Feedback." *OPSEARCH*, vol. 53, 2016, pp. 116–130. 10.1007/s12597-015-0218-5.
- XV. Saini, A., D. Gupta, and A. K. Tripathi. "Analytical Study of Two Serial Channels with Priority and Reneging." *International Journal on Recent and Innovation Trends in Computing and Communication*, vol. 11, no. 6, 2023, pp. 89–93. 10.17762/ijritcc.v11i6.7237.
- XVI. Saini, V., D. Gupta, and A. K. Tripathi. "Analysis of Feedback Queue System Comprised of Two Serial Servers with Impatient Customers." *Aryabhatta Journal of Mathematics and Informatics*, vol. 14, no. 2, 2022, pp. 145–152.

- XVII. Sangeeta, S., M. Singh, and D. Gupta. "Analysis of Serial Queues with Discouragement, Reneging and Feedback." AIP Conference Proceedings, vol. 2735, no. 1, 2023. 10.1063/12.0016682.
- XVIII. Shree, V., S. Upadhyaya, and R. Kulshrestha. "Cost Scrutiny of Discrete-Time Priority Queue with Cluster Arrival and Bernoulli Feedback." OPSEARCH, 2024, pp. 1–34. 10.1007/s12597-024-00742-8.
- XIX. Singh, T. P. "Steady State Analysis of Heterogeneous Feedback Queue Model." Proceedings of International Conference on Intelligence System and Network (ISTK), 2011.
- XX. Singh, T. P., and A. Tyagi. "Cost Analysis of a Queue System with Impatient Customer." Aryabhatta Journal of Mathematics and Informatics, vol. 6, no. 2, 2013, pp. 309–312.
- XXI. Xu, J., and L. Liu. "Analysis of a Two-Stage Tandem Queuing System with Priority and Clearing Service in the Second Stage." Mathematics, vol. 12, no. 10, 2024, p. 1500. 10.3390/math121015.
- XXII. Zadeh, A. B. "A Batch Arrival Multi-Phase Queueing System with Random Feedback in Service and Single Vacation Policy." OPSEARCH, vol. 52, 2015, pp. 617–630. 10.1007/s12597-015-0206-9

