

Analysis of Unreliable Bulk Queue with balking and Bernoulli Vacation Schedule

Binay Kumar^{1**}, Dinesh Bhaskar²

^{1*} Department of Mathematics, Patna University, Patna, India.
<https://orcid.org/0000-0002-8895-8672>

² Department of Mathematics, Patna University, Patna

Abstract: In this paper, we examine a queueing model where units arrive in randomly sized groups. The server processes all arriving units through a two-phase service, completing one phase before proceeding to the next. After completing both phases, the server may take a vacation of random duration with probability p . Additionally, the server is susceptible to random failures during any phase of service. Upon failure, it is sent for repair, after which it resumes operation as good as new. The arriving units exhibit impatience and may balk from the system without receiving service. We employ the supplementary variable technique to derive explicit expressions for various performance measures. A numerical illustration is provided to analyze the sensitivity of these performance measures.

1. Introduction

In recent years, stochastic queueing models with server vacations have gained significant attention for analyzing congestion issues and modeling various systems, including flexible manufacturing, production, and telecommunications.

In many real-world scenarios, servers may experience unexpected failures while processing units require immediate repair. Additionally, system breakdowns and repair periods can lead to fluctuations in arrival rates. This phenomenon is commonly observed in call centers, where sudden service interruptions due to system failures can cause waiting customers to leave without being served, ultimately impacting overall performance and efficiency. Similar disruptions occur in production and flexible manufacturing systems, where server breakdowns halt operations, increasing system load and affecting productivity.

Time and cost are critical factors that significantly impact the efficiency and progress of manufacturing and production systems. These systems often encounter challenges in delivering optimal output while ensuring smooth operations. In industrial settings, customer arrival rates can fluctuate due to factors such as seasonal demand variations or unexpected machine failures.

In such situations, customers waiting for service may either leave the system or switch to an alternative provider that offers faster and more cost-effective solutions. This phenomenon is evident in various sectors, including supply chain management and transportation systems, where delays caused by traffic congestion or logistic disruptions can affect overall efficiency.

A similar issue arises in airline operations. For instance, passengers waiting at an airport may experience delays due to sudden technical issues with an aircraft. If the delay is prolonged, some passengers may opt to book an alternative flight with a different airline to reach their destination on time. Such instances not only impact customer satisfaction but also affect the operational efficiency of the airline.

2. Review of literature

Customer impatience in queueing systems is a critical modeling feature with applications in call centers, network services, and customer-facing operations. It has been addressed through a wide variety of models and analytical techniques. Perel and Yechiali [1] analyze an M/M/c queue in a two-phase Markovian environment, where customer impatience during the slower service phase significantly degrades performance. Wang et al. [2] extend this analysis to the machine repair setting, incorporating balking behavior and variable server configurations to optimize service cost, bridging queueing theory with operations management. Selvaraju and Goswami [3] investigate M/M/1 queues with working vacations, offering closed-form results for both single and multiple vacation schemes. Transient dynamics are tackled by Ammar [4], who uses explicit methods to study a two-server heterogeneous queue with impatience. Kim and Kim [5] contribute a model with Markov-modulated service rates and age-based impatience, providing transient performance distributions suitable for dynamic environments. Ammar [6] also

refines transient analysis for an M/M/1 vacation queue with a waiting server and impatient customers, focusing on adaptive service policies. From a design perspective, Fu-Min et al. [7] develop a retrial model involving unreliable servers, customer feedback, and impatience, optimized using hybrid heuristics. Morozov et al. [8] explore multiclass retrial queues with balking and constant retrial rates, employing regenerative and matrix-analytic approaches validated by simulation. Finally, Chai et al. [9] study a many-to-many matching system incorporating customer impatience, applicable to platforms such as ride-sharing services.

Vacation policies model service interruptions due to rest, maintenance, or managerial decisions. Their effect on performance has been explored under various probabilistic and dynamic frameworks. Maraghi et al. [10] propose a model with batch arrivals, optional second services, and Bernoulli vacations, illustrating how stochastic scheduling balances stability and interruption risk. Rajadurai et al. [11] examine retrial queues with Bernoulli vacations and orbital search behavior, emphasizing adaptive control under heavy traffic. Liu and Wang [12] investigate how Bernoulli vacations affect equilibrium joining strategies in Markovian queues. Bouchentouf et al. [13] consider working vacations with reneging and customer retention under Bernoulli interruptions, demonstrating how partial service during vacations reduces abandonment. Vacation interruptions are further studied by Laxmi and Seleshi [14], who introduce changeover times and optimize batch-service systems using genetic algorithms. A transient analysis of such models is provided by Vijayashree and Janani [15]. Xu et al. [16] developed a queueing-inventory model with delayed Bernoulli vacations, where vacation initiation depends on inventory levels. Singh et al. [17] incorporate state-dependent arrivals and fixed-duration vacations, tuning policies to match traffic variability. Jain and Bhagat [18] unify vacations, retrials, and delayed repairs in a comprehensive model that manages multiple uncertainties simultaneously.

Mohan Kumar and Siva Kumar [19] apply metaheuristic optimization to an M/M/1 system with multiple working vacations, unreliable servers, and feedback. Mehandiratta and Verma [20] present a model with probabilistic reneging and feedback during working vacations, using analytical tools to assess sensitivity. Finally, Jain and Bhagat [21] study

a two-phase system with finite capacity and hybrid vacation schemes, revealing performance trade-offs in complex operational settings.

Queueing systems with unreliable servers model failure-prone environments common in telecommunications, healthcare, and computing. Kumar and Arumuganathan [22] studied an MX/G/1 retrial queue with active breakdowns and two-phase repairs, incorporating both patient and impatient behavior during server failures. Choudhury and Medhi [23], while focusing on multiserver systems with balking and reneging, provide complementary insights relevant to unreliable environments. Jain et al. [24] introduce an unreliable M/M/2/K queue controlled via an (N,F)-policy, including startup delays and optional repairs, and analyze transient states using matrix methods. Bhagat and Jain [25] examine a retrial queue with setup times before repair, optional second services, and unreliable servers, using generating functions for performance analysis. Jain and Gupta [26] consider redundant systems with warm standbys and switching failures, incorporating fuzzy inference into reliability evaluations. Ayyappan and Karpagam [27] investigate a queue with an unreliable main server and non-Markovian bulk service, supplemented by a standby server during breakdowns. Singh et al. [28] analyze a bulk arrival retrial queue influenced by negative customers that trigger server unreliability, modeling delayed repairs and additional service options. Saravanan et al. [29] studied a Markovian multi-server retrial system with synchronized vacation-based maintenance and discouraged customers, using QBD and matrix-geometric methods. Lastly, Kumar et al. [30] propose a healthcare queue with a single unreliable server and retrials, distinguishing urgent and non-urgent patients. By applying genetic algorithms and particle swarm optimization, the model provides cost-efficient solutions for overloaded systems.

In many practical queueing scenarios, the presence of balking, optional server vacations and server unreliability motivates the development of queueing models that incorporate batch arrivals and account for both regular and optional service periods under server unreliability.

This study is organized into several key sections. Section 3 introduces the model in detail, outlining the underlying assumptions and defining the necessary notations. In Section 4, the fundamental equations governing the system are developed using supplementary variables, along with the corresponding boundary conditions. Section 5 is dedicated to the analytical investigation of the model, focusing on both transient and steady-state behaviors of the queue length distribution through the application of Laplace transforms and probability-generating functions. Section 6 derives important performance metrics of the system. Section 7 provides numerical results and sensitivity analyses, evaluating how changes in system parameters affect various performance metrics. Graphs are used to illustrate the impact of different parameters on the average queue length, while tables present the effects of arrival and service rates on both the average number of units in the queue and the average waiting time. Finally, Section 8 presents the conclusions, highlighting the novel features and outlining the future scope of the proposed model.

3. Model Description

In many real-world scenarios, the flow of entities requiring service is influenced by the efficiency of the service facility. When the system experiences downtime due to maintenance or unexpected failures, the arrival process may be disrupted. For instance, in smart traffic management systems, if traffic signals malfunction or road maintenance causes lane closures, vehicles may be forced to reroute, leading to congestion and delays.

Similarly, in online customer service platforms, if response times are slow due to high server load or temporary outages, users may abandon their requests and seek alternatives. Such situations also arise in ticket booking systems for concerts or sports events, where slow processing or website crashes can lead customers to explore other options. By incorporating the balking nature of units, this study provides a more practical approach, demonstrating the model's relevance across various real-life applications.

To study this type of queuing problem, we examine a single-server system featuring an unreliable server, under the assumption that all internal stochastic processes

function independently. The system receives batch arrivals that follow a Poisson process with specified rates $\lambda, \lambda\varepsilon_1, \lambda\varepsilon_2, \lambda\varepsilon_3$ depending on whether the server is idle, busy, on vacation, or experiencing a breakdown respectively, where $\varepsilon_1, \varepsilon_2$ and ε_3 denote the joining probability during busy state, vacation state or repair state. The number $n (\geq 0)$ indicates the number of customer in queue. The server is subject to random failures, occurring at a rate δ , and is subsequently repaired at a rate γ ; both the operational and repair durations are modeled as exponentially distributed random variables. Let $\alpha_i(v)dv$ denote the conditional probability that the i^{th} stage of service is completed within the interval $(v, v+dv)$, given that v units of time have already elapsed. Similarly, let the $\beta(v)dv$ denote conditional probability of completing a vacation within a small interval $(v, v+dv)$, given a current vacation duration v , be represented accordingly. After finishing the second phase of service for a unit, the server may take a vacation of random duration V with a certain probability p . The time-to-completion for both the service and the optional vacation is characterized using hazard rate functions.

$$\alpha_1(v)dv = \frac{dD_1(v)}{1 - D_1(v)} \quad (1)$$

$$\alpha_2(v)dv = \frac{dD_2(v)}{1 - D_2(v)} \quad (2)$$

$$\beta(v)dv = \frac{dV(v)}{1 - V(v)} \quad (3)$$

4. Governing Equations

To analyze the non-Markovian nature of the model, the supplementary variable technique is employed to transform it into a Markovian framework. This is achieved by introducing an additional random variable representing the elapsed time of a generally distributed service time. Following the probabilistic approach outlined by Cox (1955), a set of differential-difference equations is formulated for the various states of the system. To evaluate the steady-state performance measures, the Chapman-Kolmogorov equations are constructed, accounting for state-dependent arrival rates and a two-phase service mechanism, as outlined below:

$$\frac{\partial}{\partial v} L_n^{(1)}(v, t) + \frac{\partial}{\partial t} L_n^{(1)}(v, t) + (\lambda \varepsilon_1 + \alpha_1(v) + \delta) L_n^{(1)}(v, t) = \lambda \varepsilon_1 \sum_{j=1}^n c_j L_{n-j}^{(1)}(v, t); \quad n \geq 1 \quad (4)$$

$$\frac{\partial}{\partial v} L_0^{(1)}(v, t) + \frac{\partial}{\partial t} L_0^{(1)}(v, t) + (\lambda \varepsilon_1 + \alpha_1(v) + \delta) L_0^{(1)}(v, t) = 0 \quad (5)$$

$$\frac{\partial}{\partial v} L_n^{(2)}(v, t) + \frac{\partial}{\partial t} L_n^{(2)}(v, t) + (\lambda \varepsilon_1 + \alpha_2(v) + \delta) L_n^{(2)}(v, t) = \lambda \varepsilon_1 \sum_{j=1}^n c_j L_{n-j}^{(2)}(v, t); \quad n \geq 1 \quad (6)$$

$$\frac{\partial}{\partial v} L_0^{(2)}(v, t) + \frac{\partial}{\partial t} L_0^{(2)}(v, t) + (\lambda \varepsilon_1 + \alpha_2(v) + \delta) L_0^{(2)}(v, t) = 0 \quad (7)$$

$$\frac{\partial}{\partial v} M_n(v, t) + \frac{\partial}{\partial t} M_n(v, t) + (\lambda \varepsilon_2 + \beta(v)) M_n(v, t) = \lambda \varepsilon_2 \sum_{j=1}^n c_j M_{n-j}(v, t); \quad n \geq 1 \quad (8)$$

$$\frac{\partial}{\partial v} M_0(v, t) + \frac{\partial}{\partial t} M_0(v, t) + (\lambda \varepsilon_2 + \beta(v)) M_0(v, t) = 0 \quad (9)$$

$$\frac{d}{dt} Q_n(t) = -(\lambda \varepsilon_3 + \gamma) Q_n(t) + \sum_{j=1}^n \lambda \varepsilon_3 c_j Q_{n-j}(t) + \delta \int_0^\infty L_{n-1}^{(1)}(v, t) dv + \delta \int_0^\infty L_{n-1}^{(2)}(v, t) dv; \quad n \geq 1 \quad (10)$$

$$\frac{d}{dt} Q_0(t) = -(\lambda \varepsilon_3 + \gamma) Q_0(t) \quad (11)$$

$$\frac{d}{dt} I(t) = -\mathcal{U}(t) + Q_0(t) \gamma + \int_0^\infty M_0(v, t) \beta(v) dv + (1-p) \int_0^\infty L_0^{(2)}(v, t) \alpha_2(v) dv. \quad (12)$$

The boundary and initial conditions are considered as follows:

$$L_n^{(1)}(0, t) = \lambda c_{n+1} I(t) + \gamma Q_{n+1}(t) + (1-p) \int_0^\infty L_{n+1}^{(2)}(v, t) \alpha_2(v) dv + \int_0^\infty M_{n+1}(v, t) \beta(v) dv; \quad n \geq 0 \quad (13)$$

$$L_n^{(2)}(0, t) = \int_0^\infty L_n^{(1)}(v, t) \alpha_1(v) dv \quad n \geq 0 \quad (14)$$

$$M_n(0, t) = p \int_0^\infty L_n^{(2)}(v, t) \alpha_2(v) dv \quad n \geq 0 \quad (15)$$

$$I(0) = 1 \quad (16)$$

$$M_n(0) = 0, L_n^{(j)}(0) = 0; \quad n \geq 0, \quad j = 1, 2. \quad (17)$$

Taking laplace transforms of the equations (4)-(12) give

$$\frac{\partial}{\partial v} \bar{L}_n^{(1)}(v, s) + (s + \lambda \varepsilon_1 + \alpha_1(v) + \delta) \bar{L}_n^{(1)}(v, s) = \lambda \varepsilon_1 \sum_{j=1}^n c_j \bar{L}_{n-j}^{(1)}(v, s); \quad n \geq 1 \quad (18)$$

$$\frac{\partial}{\partial v} \bar{L}_0^{(1)}(v, s) + (s + \lambda \varepsilon_1 + \alpha_1(v) + \delta) \bar{L}_0^{(1)}(v, s) = 0 \quad (19)$$

$$\frac{\partial}{\partial v} \bar{L}_n^{(2)}(v, s) + (s + \lambda \varepsilon_1 + \alpha_2(v) + \delta) \bar{L}_n^{(2)}(v, s) = \lambda \varepsilon_1 \sum_{j=1}^n c_j \bar{L}_{n-j}^{(2)}(v, s); \quad n \geq 1 \quad (20)$$

$$\frac{\partial}{\partial v} \bar{L}_0^{(2)}(v, s) + (s + \lambda \varepsilon_1 + \alpha(v) + \delta) \bar{L}_0^{(2)}(v, s) = 0 \quad (21)$$

$$\frac{\partial}{\partial v} \bar{M}_n(v, s) + (s + \lambda \varepsilon_2 + \beta(v)) \bar{M}_n(v, s) = \lambda \varepsilon_2 \sum_{j=1}^n c_j \bar{M}_{n-j}(v, s); \quad n \geq 1 \quad (22)$$

$$\frac{\partial}{\partial v} \bar{M}_0(v, s) + (s + \lambda \varepsilon_2 + \beta(v)) \bar{M}_0(v, s) = 0 \quad (23)$$

$$(s + \lambda \varepsilon_3 + \gamma) \bar{Q}_n(s) = \lambda \varepsilon_3 \sum_{j=1}^n c_j \bar{Q}_{n-j}(s) + \delta \int_0^\infty \bar{L}_{n-1}^{(1)}(v, s) dv + \delta \int_0^\infty \bar{L}_{n-1}^{(2)}(v, s) dv; \quad n \geq 1 \quad (24)$$

$$(s + \lambda \varepsilon_3 + \gamma) \bar{Q}_0(s) = 0 \quad (25)$$

$$(s + \lambda) \bar{I}(s) = 1 + \gamma \bar{Q}_0(s) + \int_0^\infty \bar{M}_0(v, s) \beta(v) dv + (1 - p) \int_0^\infty \bar{L}_0^{(2)}(v, s) \alpha_2(v) dv. \quad (26)$$

By Laplace transforms of the equations (13)-(15), we have

$$\bar{L}_n^{(1)}(0, s) = \lambda c_{n+1} \bar{I}(s) + \gamma \bar{Q}_{n+1}(s) + (1 - p) \int_0^\infty \bar{L}_{n+1}^{(2)}(v, s) \alpha_2(v) dv + \int_0^\infty \bar{M}_{n+1}(v, s) \beta(v) dv; \quad n \geq 0 \quad (27)$$

$$\bar{L}_n^{(2)}(0, s) = \int_0^\infty \bar{L}_n^{(1)}(v, s) \alpha_1(v) dv; \quad n \geq 0 \quad (28)$$

$$\bar{M}_n(0, s) = p \int_0^\infty \bar{L}_n^{(2)}(v, s) \alpha_2(v) dv; \quad n \geq 0. \quad (29)$$

5. Mathematical Analysis

In the stochastic modeling of queuing systems, mathematical analysis plays a crucial role in evaluating various performance metrics under specific assumptions. Through a comprehensive performance analysis, we gain insights into the system's behavior and efficiency.

In this section, we utilize the probability-generating function technique to derive queue size distributions across different scenarios, considering the server's varying operational states. Additionally, we examine both transient and steady-state behaviors of the model at random time points.

For computation purpose, we use the following notation throughout the paper.

$$\bar{\eta}_1(s, z) = (s + \lambda \varepsilon_1(1 - X(z)) + \delta), \bar{\eta}_2(s, z) = (s + \lambda \varepsilon_2(1 - X(z))), \bar{\eta}_3(s, z) = (s + \lambda \varepsilon_3(1 - X(z)) + \gamma),$$

$$\bar{\pi}(z, s) = \bar{\eta}_1 \bar{\eta}_3 [z - p \bar{D}_1(\bar{\eta}_1) \bar{D}_2(\bar{\eta}_1) \bar{V}(\bar{\eta}_2) - (1 - p) \bar{D}_1(\bar{\eta}_1) \bar{D}_2(\bar{\eta}_1)]$$

$$- (\gamma \delta z (1 - \bar{D}_1(\bar{\eta}_1) \bar{D}_2(\bar{\eta}_1))).$$

And

$$\eta_1(z) = \lim_{s \rightarrow 0} \bar{\eta}_1(s, z), \eta_2(z) = \lim_{s \rightarrow 0} \bar{\eta}_2(s, z), \eta_3(z) = \lim_{s \rightarrow 0} \bar{\eta}_3(s, z), \pi(z) = \lim_{s \rightarrow 0} \bar{\pi}(s, z)$$

Theorem 1: The probability generating functions corresponding to the transient state at a random time, when the server is engaged in service, on vacation, or undergoing repair, are respectively expressed as follows:

$$\bar{L}^{(1)}(z, s) = \frac{\bar{\eta}_3(s, z)[1 - sI(s) + \lambda(X(z) - 1)\bar{I}(s)][1 - \bar{D}_1(\bar{\eta}_1(s, z))]}{\bar{\pi}(z, s)} \quad (30)$$

$$\bar{L}^{(2)}(z, s) = \frac{\bar{\eta}_3(s, z)[1 - sI(s) + \lambda(X(z) - 1)\bar{I}(s)]\bar{D}_1(\bar{\eta}_1(s, z))[1 - \bar{D}_2(\bar{\eta}_1(s, z))]}{\bar{\pi}(z, s)} \quad (31)$$

$$\bar{M}(z, s) = \frac{p\bar{\eta}_1(s, z)\bar{\eta}_3(s, z)[1 - sI(s) + \lambda(X(z) - 1)\bar{I}(s)]\bar{D}_1(\bar{\eta}_1(s, z))\bar{D}_2(\bar{\eta}_1(s, z))[1 - \bar{V}(\bar{\eta}_2(s, z))]}{\bar{\pi}(z, s)\eta_2(s, z)} \quad (32)$$

$$\bar{Q}(z, s) = \frac{\delta z[1 - sI(s) + \lambda(X(z) - 1)\bar{I}(s)][1 - \bar{D}_1(\bar{\eta}_1(s, z))\bar{D}_2(\bar{\eta}_1(s, z))]}{\bar{\pi}(z, s)} \quad (33)$$

Proof: For proof see Appendix-A.

Theorem 2: The steady-state probability generating functions observed at a random point in time, corresponding to the server being in service, on vacation, or under repair, are respectively represented as follows:

$$L^{(1)}(z) = \frac{(1 - \rho)\eta_3(z)[\lambda(X(z) - 1)][1 - \bar{D}_1(\eta_1(z))]}{\pi(z)} \quad (34)$$

$$L^{(2)}(z) = \frac{(1 - \rho)\eta_3(z)[\lambda(X(z) - 1)]\bar{D}_1(\eta_1(z))[1 - \bar{D}_2(\eta_1(z))]}{\pi(z)} \quad (35)$$

$$M(z) = \frac{(1 - \rho)p\eta_1(z)\eta_3(z)\bar{D}_1(\eta_1(z))\bar{D}_2(\eta_1(z))[\bar{V}(\eta_2(z)) - 1]}{\varepsilon_2\pi(z)} \quad (36)$$

$$Q(z) = \frac{(1 - \rho)\delta z[\lambda(X(z) - 1)][1 - \bar{D}_1(\eta_1(z))\bar{D}_2(\eta_1(z))]}{\pi(z)} \quad (37)$$

Proof: For proof see Appendix-B.

Theorem 3: The probability generating function of the number of units in the queue at a random epoch is given by

$$P(z) = \frac{(1-\rho)[\eta_2(z)\{\bar{D}_1(\eta_1(z))\bar{D}_1(\eta_1(z))-1\}\{\eta_2(z)+\delta z\} + p\eta_1(z)\eta_3(z)\bar{D}_1(\eta_1(z))\bar{D}_2(\eta_1(z))\{\bar{V}(\eta_2(z))-1\}]}{\varepsilon_2\pi(z)} \quad (38)$$

Proof: On adding the equations (34)-(37), we get the required equation (38).

6. Performance Measures

Performance metrics are essential for examining the behavior of the system when creating a new one or improving an existing queueing model. These steps aid in locating and controlling blocking and delays that are frequently seen in actual traffic situations. Because queueing systems are widely used in both everyday and industrial processes, system performance can differ greatly depending on the circumstances. In this section, we use the previously obtained probability generating functions to derive explicit expressions for key performance indicators, specifically the average number of units in the queue.

(i) Average number of units in the queue (system)

The equation (38) can be written as

$$P(z) = \frac{N(z)}{D(z)} \quad (39)$$

$$n(z) = (1-\rho)[\eta_2(z)\{\bar{D}_1(\eta_1(z))\bar{D}_1(\eta_1(z))-1\}\{\eta_2(z)+\delta z\} + p\eta_1(z)\eta_3(z)\bar{D}_1(\eta_1(z))\bar{D}_2(\eta_1(z))\{\bar{V}(\eta_2(z))-1\}]$$

$$D(z) = \varepsilon_2\pi(z)$$

On differentiating the numerator and denominator respectively of right-hand side of equation (39) at $z = 1$, we get

$$N'(1) = (1-\rho)\lambda E(X)\{(\delta+\gamma) + \bar{D}_1(\delta)\bar{D}_2(\delta)[p\delta\gamma E(V) - \delta - \gamma]\} \quad (40)$$

$$\begin{aligned} N''(1) = & (1-\rho)\left[2\lambda(\lambda\varepsilon_1)(E(X))^2\left\{\bar{D}_1'(\delta)\bar{D}_2(\delta) + \bar{D}_1(\delta)\bar{D}_2'(\delta)\right\}(\delta+\gamma) - p\delta\gamma E(V)\right\} \\ & + \left\{-\frac{\delta}{E(X)} - \lambda\varepsilon_3\right\} + \bar{D}_1(\delta)\bar{D}_2(\delta)\left\{\lambda\varepsilon_3 - \frac{\delta}{E(X)} - p(\varepsilon_3\delta + \varepsilon_1\gamma)\lambda E(V) + \frac{\lambda\varepsilon_2}{2}p\delta\gamma E(V^2)\right\}] \\ & + \lambda E(X^{(2)})[(\delta+\gamma) + \bar{D}_1(\delta)\bar{D}_2(\delta)(p\delta\gamma E(V) - \delta - \gamma)] \end{aligned} \quad (41)$$

$$D'(1) = -\lambda E(X)(\varepsilon_3 \delta + \varepsilon_1 \gamma) + \bar{D}_1(\delta) \bar{D}_2(\delta) \{ \delta \gamma + \lambda E(X) [(\varepsilon_3 \delta + \varepsilon_1 \gamma) - \varepsilon_2 \delta \gamma p E(V)] \} \quad (42)$$

$$D''(1) = 2(\lambda E(X))^2 \left[\varepsilon_3 \varepsilon_1 - \frac{(\varepsilon_3 \delta + \varepsilon_1 \gamma)}{\lambda E(X)} \right] + \bar{D}_1(\delta) \bar{D}_2(\delta) \left\{ p \varepsilon_2 (\varepsilon_3 \delta + \varepsilon_1 \gamma) E(V) - \varepsilon_1 \varepsilon_3 - \frac{(\varepsilon_2)^2}{2} p \delta \gamma E(V^2) \right\} \\ + \left[\bar{D}_1'(\delta) \bar{D}_2(\delta) + \bar{D}_1(\delta) \bar{D}_2'(\delta) \right] \left\{ -\varepsilon_1 (\varepsilon_3 \delta + \varepsilon_1 \gamma) + p \varepsilon_1 \varepsilon_2 \delta \gamma E(V) - \frac{\varepsilon_1 \delta \gamma}{\lambda E(X)} \right\} \\ + \lambda E(X^{(2)}) \left[-\varepsilon_3 \delta - \varepsilon_1 \gamma + \bar{D}_1(\delta) \bar{D}_2(\delta) (\varepsilon_3 \delta + \varepsilon_1 \gamma - p \varepsilon_2 \delta \gamma E(V)) \right] \quad (43)$$

By substituting the results of equations (40)-(43) in equation

$$L_q = \left[\frac{d}{dz} P(z) \right]_{z=1} = \frac{D'(1)N''(1) - D''(1)N'(1)}{(D''(1))^2} \quad (44)$$

we get the required result of the average number of units in the queue.

(ii) Average waiting time of the units in the queue

The average waiting time (W_q) of the units in the queue can be determined by using equation (2.19) with $\lambda_e = \lambda \varepsilon_1 (L^{(1)}(1) + L^{(2)}(1)) + \lambda I + \lambda \varepsilon_2 M(1) + \lambda \varepsilon_3 Q(1)$.

7. Numerical Illustration and Sensitivity Analysis

The average number of units and average waiting time that were determined analytically in the preceding part are validated numerically in this section. The batch size of the units is assumed to follow the geometric distribution with first and second moments as provided in equation in order to assist the numerical results.

$$E(X) = \frac{b}{a}, E(X^2) = \frac{b(1+b)}{a^2}; b = 1-a, 0 < a < 1. \quad (45)$$

The service time distribution is assumed to follow exponential distribution with parameter $\mu_i (i = 1, 2)$ s and the server in optional vacation state with parameters ν is also exponentially distributed. The first and second moments are defined as

$$E(B_i) = \frac{1}{\mu_i}, E(B_i^2) = \frac{2}{\mu_i^2}; i = 1, 2. \quad (46)$$

Further, $\bar{B}_i(\alpha)$ and $\bar{B}_i'(\alpha)$ can be obtained by using

$$\bar{B}_i(\alpha) = \frac{\mu_i}{(\alpha + \mu_i)} \quad \text{and} \quad \bar{B}_i'(\alpha) = \frac{-\mu_i}{(\alpha + \mu_i)^2}, \quad \text{for } i = 1, 2. \quad (47)$$

Also

$$E(V) = \frac{1}{\nu} \quad \text{and} \quad E(V^2) = \frac{2}{\nu^2}. \quad (48)$$

Table 1: Effects of λ and μ on L_q and W_q

λ	$\mu = 6$		$\mu = 6.2$		$\mu = 6.4$	
	L_q	W_q	L_q	W_q	L_q	W_q
0.38	19.31	19.79	16.05	16.36	13.74	13.95
0.39	22.67	22.72	18.39	18.34	15.49	15.38
0.4	27.15	26.64	21.35	20.84	17.62	17.12
0.41	33.45	32.13	25.22	24.10	20.27	19.28
0.42	42.94	40.41	30.47	28.53	23.66	22.05
0.43	58.85	54.31	38.03	34.90	28.14	25.71
0.44	91.09	82.45	49.82	44.85	34.37	30.79

Table 2: Effects of δ and γ on L_q

δ	$\gamma = 6.4$		$\gamma = 6.6$		$\gamma = 6.8$	
	L_q	W_q	L_q	W_q	L_q	W_q
3	21.35	20.84	20.44	19.90	19.63	19.08
3.2	25.80	25.35	24.45	23.97	23.29	22.78
3.4	32.26	31.90	30.15	29.74	28.37	27.92
3.6	42.47	42.28	38.83	38.56	35.89	35.55
3.8	61.06	61.21	53.73	53.72	48.20	48.05

Table 3: Effects of p and δ on L_q

p	$\delta = 3$		$\delta = 3.2$		$\delta = 3.4$	
	L_q	W_q	L_q	W_q	L_q	W_q
0.1	11.28	10.74	12.74	12.20	14.51	14.00
0.3	14.55	13.99	16.73	16.19	19.52	19.02
0.5	19.63	19.08	23.29	22.78	28.37	27.92
0.7	28.68	28.16	36.14	35.71	48.27	47.99
0.9	49.36	48.96	72.85	72.71	135.10	135.67

Table 4: Effects of δ and ν on L_q

δ	$\nu = 4$		$\nu = 4.5$		$\nu = 5$	
	L_q	W_q	L_q	W_q	L_q	W_q
3.0	19.63	19.08	17.94	17.39	16.75	16.20
3.2	23.29	22.78	21.06	20.54	19.52	19.00
3.4	28.37	27.92	25.27	24.80	23.20	22.72
3.6	35.89	35.55	31.29	30.91	28.32	27.91
3.8	48.20	48.05	40.58	40.35	35.93	35.65
4.0	71.97	72.23	56.81	56.86	48.47	48.40

Table 1 presents a detailed analysis of the impact of arrival and service rates on the average queue length and waiting time. The results indicate that an increase in the arrival rate leads to a corresponding rise in both metrics, reflecting increased system workload and congestion. In contrast, a higher service rate results in a reduction in average queue length and waiting time, highlighting the role of enhanced service capacity in alleviating congestion and improving overall system performance. Table 2 examines the influence of failure and repair rates on these performance indicators. It is observed that higher failure rates lead to increases in both average queue length and waiting time, due to more frequent service interruptions. Conversely, improvements in the repair rate lead to a consistent decline in both metrics, emphasizing the importance of efficient maintenance in enhancing system reliability and throughput. Table 3 explores the combined effects of the probability parameter p and the failure rate on system performance. The findings reveal that both average queue length and waiting time increase with a rise in the failure rate, while higher values of p further exacerbate this trend, indicating that both factors jointly contribute to system congestion and reduced efficiency. Finally, Table 4 investigates the interplay between the failure rate (δ) and the vacation rate (ν) in determining queueing performance. As the vacation rate increases, both average queue length and waiting time exhibit a declining trend, suggesting improved server availability and efficiency. However, an increase in the failure rate causes a sharp escalation in both metrics, highlighting the significant adverse impact of frequent failures on overall system behavior.

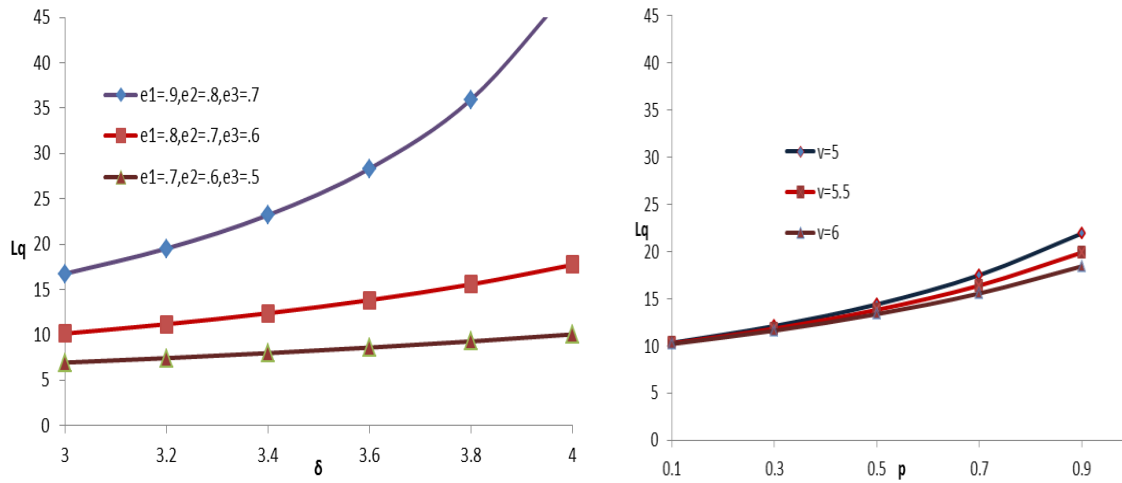


Fig. 1: L_q vs. δ for different balking rates **Fig. 2:** L_q vs. p for different values of v

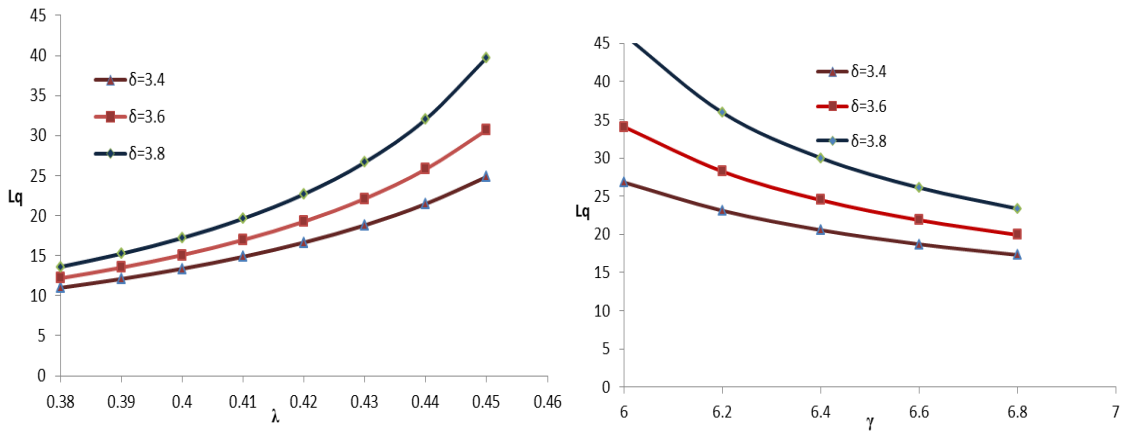


Fig. 3: L_q vs. λ for different values of δ **Fig. 4:** L_q vs. γ for different values of δ

Figure 1 illustrates the variation in average queue length with respect to changes in the failure rate (δ) across different balking rates. As the joining rate decreases—corresponding to an increase in the balking rate—the average queue length declines, primarily because more units leave the system without receiving service. Furthermore, an

increase in the failure rate leads to a longer queue, with this effect being more pronounced under higher joining rates. Figure 2 demonstrates the effect of the parameter p on the average queue length under varying vacation rates. As the vacation rate increases, the average queue length tends to decrease, as the server becomes available more frequently. In contrast, an increase in p results in a longer queue since the server takes vacations more often, thereby reducing service availability. Figure 3 highlights the impact of changes in the arrival rate on average queue length under different failure rates. A clear upward trend is observed in queue length with increasing arrival and failure rates, indicating heightened system congestion. Finally, Figure 4 presents the variation in average queue length with respect to changes in the repair rate under various failure rates. As the repair rate increases, the average queue length decreases due to improved recovery of failed components, whereas higher failure rates lead to longer queues as a result of more frequent interruptions in service.

8. Conclusion

In this paper, we have examined a stochastic queueing model featuring two stages of heterogeneous service, each with arbitrarily distributed service times. The model incorporates an unreliable server that may optionally take a vacation upon completing both service stages for a unit. This framework offers greater flexibility and is particularly relevant for manufacturing and production environments where machines require scheduled downtime or maintenance breaks of fixed duration. The analytical findings provide key performance metrics, which can be utilized in cost-based optimization to enhance system efficiency. By accounting for service interruptions, the model effectively captures the dynamics of various industrial and real-world scenarios. It serves as a valuable tool for system designers and queueing theorists to assess and manage delays and blocking issues in practical applications.

Reference:

- [1] N. Perel and U. Yechiali, Queues with slow servers and impatient customers, *Eur. J. Oper. Res.* **201** (2010), no. 1, 247–258. <https://doi.org/10.1016/j.ejor.2009.02.024>

- [2] K.-H. Wang, Y.-C. Liou, and D.-Y. Yang, Cost optimization and sensitivity analysis of the machine repair problem with variable servers and balking, *Procedia Soc. Behav. Sci.* **25** (2011), 178–188. <https://doi.org/10.1016/j.sbspro.2011.10.539>
- [3] N. Selvaraju and C. Goswami, Impatient customers in an M/M/1 queue with single and multiple working vacations, *Comput. Ind. Eng.* **65** (2013), no. 2, 207–215. <https://doi.org/10.1016/j.cie.2013.02.016>
- [4] S. I. Ammar, Transient analysis of a two-heterogeneous servers queue with impatient behavior, *J. Egypt. Math. Soc.* **22** (2014), no. 1, 90–95. <https://doi.org/10.1016/j.joems.2013.05.002>
- [5] B. Kim and J. Kim, A single server queue with Markov modulated service rates and impatient customers, *Perform. Eval.* **83–84** (2015), 1–15. <https://doi.org/10.1016/j.peva.2014.11.002>
- [6] S. I. Ammar, Transient solution of an M/M/1 vacation queue with a waiting server and impatient customers, *J. Egypt. Math. Soc.* **25** (2017), no. 3, 337–342. <https://doi.org/10.1016/j.joems.2016.09.002>
- [7] C. Fu-Min, L. Tzu-Hsin, and K. Jau-Chuan, On an unreliable-server retrial queue with customer feedback and impatience, *Appl. Math. Modelling* **55** (2018), 171–182. <https://doi.org/10.1016/j.apm.2017.10.025>
- [8] E. Morozov, A. Rumyantsev, S. Dey, and T. G. Deepak, Performance analysis and stability of multiclass orbit queue with constant retrial rates and balking, *Perform. Eval.* **134** (2019), 102005. <https://doi.org/10.1016/j.peva.2019.102005>
- [9] X. Chai et al., On a many-to-many matched queueing system with flexible matching mechanism and impatient customers, *J. Comput. Appl. Math.* **416** (2022), 114573. <https://doi.org/10.1016/j.cam.2022.114573>
- [10] H. Maraghi, M. Khedmati, and M. A. B. Jahromi, Analyzing an M[X]/G/1 queue with second optional service, Bernoulli schedule vacation and random breakdowns, *RAIRO - Operations Research* **53** (2019), no. 4, 1309–1333.
- [11] R. Rajadurai, V. N. Janaki, and R. Vanjinathan, Performance analysis of a retrial queue with orbital search, Bernoulli vacation, balking and feedback, *Yugoslav J. Oper. Res.* **29** (2019), no. 1, 79–98.
- [12] Z. Liu and K. Wang, Equilibrium and social optimal strategies in observable queues with Bernoulli schedule vacation, *Comput. Ind. Eng.* **90** (2015), 122–129.

- [13] A. Bouchentouf, A. L. A. Said, and R. Errouissi, Queueing model with Bernoulli schedule vacation interruption, reneging, and retention of reneged customers, *J. Math.* **2021**, Article ID 8810221.
- [14] M. Laxmi and Y. Seleshi, Batch service queue with Bernoulli vacation interruption and changeover time, *Int. J. Math. Oper. Res.* **23** (2023), no. 1, 94–115.
- [15] T. Vijayashree and R. Janani, Transient analysis of $M[x]/M/1$ queue with Bernoulli vacation interruption, *Int. J. Sci. Res. Math. Stat. Sci.* **7** (2020), no. 4, 81–91.
- [16] X. Xu, W. Yue, and J. Cao, A queueing-inventory model with delayed Bernoulli vacation under N-policy, *Mathematics* **10** (2022), no. 21, Article ID 4109.
- [17] S. Singh, S. Saini, and R. Kumar, Analysis of $M[x]/M/1$ queue with state-dependent arrivals, fixed vacation and multiple working vacations, *OPSEARCH* **60** (2023), no. 2, 460–475.
- [18] M. Jain and A. Bhagat, A single server queue with Bernoulli vacation interruption, unreliable server and delayed repair, *Int. J. Appl. Manag. Sci.* **14** (2022), no. 2, 157–180.
- [19] G. Mohankumar and V. Sivakumar, Optimal control of an $M/M/1$ queueing system with multiple working vacations, unreliable server and feedback, *Int. J. Math. Oper. Res.* **22** (2022), no. 2, 246–268.
- [20] R. Mehandiratta and A. Verma, Performance modeling of a single server feedback queue with Bernoulli working vacations and probabilistic reneging, *Int. J. Math. Eng. Manag. Sci.* **7** (2022), no. 3, 417–431.
- [21] M. Jain and A. Bhagat, Transient analysis of an $M/M/1/N$ feedback queue with working vacation, vacation and customer impatience, *Yugoslav J. Oper. Res.* **31** (2021), no. 2, 219–241.
- [22] M. S. Kumar and R. Arumuganathan, An $MX/G/1$ retrial queue with two-phase service subject to active server breakdowns and two types of repair, *Int. J. Oper. Res.* **8** (2010), no. 3, 261–291. <https://doi.org/10.1504/IJOR.2010.033540>
- [23] A. Choudhury and P. Medhi, Balking and reneging in multiserver Markovian queueing system, *Int. J. Math. Oper. Res.* **3** (2011), no. 4, 377–394. <https://doi.org/10.1504/IJMOR.2011.040874>

- [24] M. Jain, G. C. Sharma, and R. Sharma, Optimal control of (N, F) policy for unreliable server queue with multi-optional phase repair and start-up, *Int. J. Math. Oper. Res.* **4** (2012), no. 2, 152–174. <https://doi.org/10.1504/IJMOR.2012.046375>
- [25] A. Bhagat and M. Jain, Unreliable MX/G/1 retrial queue with multi-optional services and impatient customers, *Int. J. Oper. Res.* **17** (2013), no. 2, 248–273.
- [26] M. Jain and R. Gupta, N-policy for redundant repairable system with multiple types of warm standbys with switching failure and vacation, *Int. J. Math. Oper. Res.* **13** (2018), no. 4, 419–449. <https://doi.org/10.1504/IJMOR.2018.095484>
- [27] G. Ayyappan and S. Karpagam, An MX/G(a,b)/1 queueing system with server breakdown and repair, stand-by server and single vacation, *Int. J. Math. Oper. Res.* **14** (2019), no. 2, 221–235. <https://doi.org/10.1504/IJMOR.2019.097756>
- [28] C. J. Singh, S. Kaur, and M. Jain, Unreliable server retrial G-queue with bulk arrival, optional additional service and delayed repair, *Int. J. Oper. Res.* **38** (2020), no. 1, 82–111. <https://doi.org/10.1504/IJOR.2020.106362>
- [29] V. Saravanan, V. Poongothai, and P. Godhandaraman, Performance analysis of a multi server retrial queueing system with unreliable server, discouragement and vacation model, *Math. Comput. Simul.* **214** (2023), 204–226. <https://doi.org/10.1016/j.matcom.2023.07.008>
- [30] K. Kumar, S. Garg, and S. Sharma, Cost optimization for a retrial repairable queueing healthcare system with working vacation and unreliability using the genetic algorithm and particle swarm optimization, *Oper. Res. Data Anal. Logist.* **45** (2025), Art. no. 200473. <https://doi.org/10.1016/j.ordal.2025.200473>

Appendix-A

Proof of Theorem 1:

By multiplying equations (18) and (19) with appropriate powers of z and summing over all possible values of the n , we obtain

$$\frac{\partial}{\partial v} \bar{L}^{(1)}(v, z, s) + (\bar{\eta}_1(s, z) + \alpha_1(v)) \bar{L}^{(1)}(v, z, s) = 0. \quad (\text{A.1})$$

Also from equations (20)-(25), we have

$$\frac{\partial}{\partial v} \bar{L}^{(2)}(v, z, s) + (\bar{\eta}_1(s, z) + \alpha_2(v)) \bar{L}^{(2)}(v, z, s) = 0 \quad (\text{A.2})$$

$$\frac{\partial}{\partial v} \bar{M}(v, z, s) + (\bar{\eta}_2(s, z) + \beta(v)) \bar{M}(v, z, s) = 0 \quad (\text{A.3})$$

$$\bar{\eta}_3(s, z) \bar{Q}(z, s) = \delta z \left[\int_0^\infty \bar{L}^{(1)}(v, z, s) dv + \int_0^\infty \bar{L}^{(2)}(v, z, s) dv \right]. \quad (\text{A.4})$$

Applying the same treatment for boundary condition (27), we get

$$\begin{aligned} z \bar{L}^{(1)}(0, z, s) &= \lambda(X(z) - 1) \bar{I}(s) + (1 - s \bar{I}(s)) + \gamma \bar{Q}(z, s) \\ &\quad + \int_0^\infty \bar{M}(v, z, s) \beta(v) dv + (1 - p) \int_0^\infty \bar{L}^{(2)}(v, z, s) \alpha_2(v) dv. \end{aligned} \quad (\text{A.5})$$

On same treatment, the equations (28) and (29) give

$$\bar{L}^{(2)}(0, z, s) = \int_0^\infty \bar{L}^{(1)}(v, z, s) \alpha_1(v) dv \quad (\text{A.6})$$

$$\bar{M}(0, z, s) = p \int_0^\infty \bar{L}^{(2)}(v, z, s) \alpha_2(v) dv. \quad (\text{A.7})$$

On Solving the equation (A.1) gives

$$\bar{L}^{(1)}(v, z, s) = \bar{L}^{(1)}(0, z, s) \exp\{-\bar{\eta}_1(s, z)v\} [1 - D_1(v)] \quad (\text{A.8})$$

Integrate the equation (A.8) with respect to v with limit 0 to ∞ , we have

$$\bar{L}^{(1)}(z, s) = \int_0^\infty \bar{L}^{(1)}(v, z, s) dv = \bar{L}^{(1)}(0, z, s) \left[\frac{1 - \bar{D}_1(\bar{\eta}_1(s, z))}{\bar{\eta}_1(s, z)} \right] \quad (\text{A.9})$$

Multiplying equation (A.8) by $\alpha_1(v)$ and integrating over v with limit 0 to ∞ , we have

$$\int_0^\infty \bar{L}^{(1)}(v, z, s) \alpha_1(v) dv = \bar{L}^{(1)}(0, z, s) \bar{D}_1(\bar{\eta}_1(s, z)). \quad (\text{A.10})$$

On solving the equations (A.2)-(A.3), we get

$$\bar{L}^{(2)}(v, z, s) = \bar{L}^{(2)}(0, z, s) \exp\{-\bar{\eta}_1(s, z)v\} [1 - D_2(v)] \quad (\text{A.11})$$

$$\bar{M}(v, z, s) = \bar{M}(0, z, s) \exp\{-\bar{\eta}_2(s, z)v\} [1 - V(v)] \quad (\text{A.12})$$

Integrating the equations (A.11) and (A.12) with limit 0 to ∞ gives.

$$\bar{L}^{(2)}(z, s) = \bar{L}^{(2)}(0, z, s) \left[\frac{1 - \bar{D}_2(\bar{\eta}_1(s, z))}{\bar{\eta}_1(s, z)} \right] \quad (\text{A.13})$$

$$\bar{M}(z, s) = \bar{M}(0, z, s) \left[\frac{1 - \bar{V}(\bar{\eta}_2(s, z))}{\bar{\eta}_2(s, z)} \right] \quad (\text{A.14})$$

Multiplying the equation (A.11) by $\alpha_2(v)$ and integrating with limit 0 to ∞ , we get

$$\int_0^\infty \bar{L}^{(2)}(v, z, s) \alpha_2(v) dv = \bar{L}^{(2)}(0, z, s) \bar{D}_2(\bar{\eta}_1(s, z)). \quad (\text{A.15})$$

Multiplying both sides of the equation (A.12) by $\beta(v)$ and integrating with limits 0 to ∞ , we get

$$\int_0^\infty \bar{M}(v, z, s) \beta(v) dv = \bar{M}(0, z, s) \bar{V}(\bar{\eta}_2(s, z)). \quad (\text{A.16})$$

Using equations (A.6), (A.10) and (A.15), we can write equation (A.7) as

$$\bar{M}(0, z, s) = p \bar{L}^{(1)}(0, z, s) \bar{D}_1(\bar{\eta}_1(s, z)) \bar{D}_2(\bar{\eta}_1(s, z)). \quad (\text{A.17})$$

Using equation (A.17) in equation (A.14), we have

$$\bar{M}(z, s) = p \bar{L}^{(1)}(0, z, s) \bar{D}_1(\bar{\eta}_1(s, z)) \bar{D}_2(\bar{\eta}_1(s, z)) \left[\frac{1 - \bar{V}(\bar{\eta}_2(s, z))}{\bar{\eta}_2(s, z)} \right]. \quad (\text{A.18})$$

By using equation (A.17), equation (A.16) becomes

$$\int_0^\infty \bar{M}(v, z, s) \beta(v) dv = p \bar{L}^{(1)}(0, z, s) \bar{D}_1(\bar{\eta}_1(s, z)) \bar{D}_2(\bar{\eta}_1(s, z)) \bar{V}(\bar{\eta}_2(s, z)). \quad (\text{A.19})$$

Now using (A.10), equation (A.6) reduces to

$$\bar{L}^{(2)}(0, z, s) = \bar{L}^{(1)}(0, z, s) \bar{D}_1(\bar{\eta}_1(s, z)). \quad (\text{A.20})$$

Using equations (A.13), (A.15) and (A.20), equation (A.4) becomes

$$\bar{Q}(z, s) = \bar{\alpha} \bar{L}^{(1)}(0, z, s) \frac{[1 - \bar{D}_1(\bar{\eta}_1(s, z)) \bar{D}_2(\bar{\eta}_1(s, z))]}{\bar{\eta}_1(s, z) \bar{\eta}_3(s, z)}. \quad (\text{A.21})$$

By equations (A.5) and (A.21), we get

$$\bar{L}^{(1)}(0, z, s) = \frac{\bar{\eta}_1 \bar{\eta}_3 [1 - sI(s) + \lambda(X(z) - 1)\bar{I}(s)]}{\pi(z, s)}. \quad (\text{A.22})$$

Now substituting the value of $\bar{L}^{(i)}(0, z, s)$ for $i=1,2$ from equations (A.20) and (A.22) into equations (A.9), (A.13), (A.18), and (A.21), we get the required result.

Appendix-B

Proof of Theorem 2:

Multiply equations (34)-(37) by s , taking the limit as $s \rightarrow 0$ and using Tauberian property, $\lim_{s \rightarrow 0} s \bar{f}(s) = \lim_{t \rightarrow \infty} f(t)$ we obtain equations

$$L^{(1)}(z) = \frac{\eta_3(z) [\lambda(X(z) - 1) [1 - \bar{D}_1(\eta_1(z))] I]}{\pi(z)} \quad (\text{B.1})$$

$$L^{(2)}(z) = \frac{\eta_3(z)[\lambda(X(z)-1)][\bar{D}_1(\eta_1(z))][1-\bar{D}_2(\eta_1(z))I]}{\pi(z)} \quad (\text{B.2})$$

$$M(z) = \frac{p\eta_1(z)\eta_2(z)\bar{D}_1(\eta_1(z))\bar{D}_2(\eta_1(z))[\bar{V}(\eta_2(z))-1]I}{\varepsilon_2\pi(z)} \quad (\text{B.3})$$

$$Q(z) = \frac{\delta[\lambda(X(z)-1)]\{1-\bar{D}_1(\eta_1(z))\bar{D}_2(\eta_1(z))\}I}{\pi(z)}. \quad (\text{B.4})$$

Taking limit $z \rightarrow 1$ in equations (B.1)-(B.4) and using the normalizing condition as

$$L^{(1)}(1) + L^{(2)}(1) + M(1) + Q(1) + I = 1. \quad (\text{B.5})$$

We get the required value of I as

$$I = 1 - \rho \quad (\text{B.6})$$

where

$$\rho = \frac{\lambda E(X)[(\delta + \gamma)\{1 - \bar{D}_1(\delta)\bar{D}_2(\delta)\} + p\delta\gamma\bar{D}_1(\delta)\bar{D}_2(\delta)E(V)]}{\delta\gamma\bar{D}_1(\delta)\bar{D}_2(\delta) + \lambda E(X)[(1 - \bar{D}_1(\delta)\bar{D}_2(\delta))\{(1 - \varepsilon_3)\delta + (1 - \varepsilon_1)\gamma\} + p\delta\gamma\bar{D}_1(\delta)\bar{D}_2(\delta)(1 - \varepsilon_2)E(V)]}. \quad (\text{B.7})$$

Using the equation (B.6) in equations (B.1)-(B.4) we get the required result.