

Application of fuzzy queuing models in a range of practical contexts.

Manisha Yadav¹, Dr. Naveen Kumar², Dr. Sunil Kumar³

¹Research Scholar, Dept. of Mathematics, BMU, Asthal Bohar, Rohtak

²Professor, Dept. of Mathematics, BMU, Asthal Bohar, Rohtak

³Associate Professor, Dept. of Mathematics, Govt. College for Women, Rewari

Abstract:

People frequently encounter lines in a variety of settings in today's fast-paced culture, including traffic lights, hospitals, and customer service counters. Queue time can result in ineffective waiting, which affects system efficiency and consumer pleasure. By evaluating performance measures like average wait times and system states, queuing theory—a mathematical study of waiting lines—is used to examine and optimize these procedures. With an emphasis on Indian railway ticket reservation systems and hospital bed allocation, this study investigates the applicability of queuing theory in a variety of fields. By figuring out the ideal number of servers, we use the M/M/S queuing model to investigate ways to balance client wait times and service expenses. Furthermore, we examine patterns of bulk arrivals and services, with a focus on the transportation and healthcare sectors, where efficient queuing management can greatly enhance the quality of services. In order to improve system performance and wait times and, ultimately, the user experience in public service systems, our results highlight the significance of queuing theory as a tool.

Keywords: public utilities, hospital bed allocation, waiting time reduction, bulk queues, M/M/S model, queuing theory, and service optimization.

Introduction:

Even while effective time management is valued in today's fast-paced culture, standing in line is still an inevitable part of daily life. Many governmental and private sectors, such as banking, retail, healthcare, and transportation, involve lines. Customer discontent and inefficiencies in service delivery are frequently caused by the waiting times in these lines, also referred to as

queues. As a result, efficient queuing systems are now essential for reducing wait times as well as enhancing overall service quality and operational effectiveness.

Organizations can better manage service delivery and maximize resources by using queuing theory, which offers a methodical way to researching waiting lines. Queuing theory was first developed in the contexts of operations research and probability, but it is currently widely used in a variety of industries, including industrial engineering, healthcare, and telecommunications. Queuing models, for example, have been useful in evaluating patient flow in ERs, streamlining call center operations, and enhancing the effectiveness of public transit systems. Organizations can improve their decision-making about resource allocation by using queuing theory to identify key performance measures including average waiting time, service utilization, and system capacity (Gross & Harris, 2008; Bhat, 2015).

This study aims to investigate the applications of queuing theory in two main domains: Indian railway ticket reservation systems and hospital bed allocation. These situations illustrate important public services where efficient queuing management can result in major advantages including shorter wait times, better service, and happier customers. In particular, we assess how optimal server allocation affects wait time reduction and service cost management using the M/M/S queuing model. The goal of the research is to provide practical solutions that can increase productivity and advance public welfare by comprehending queuing behavior in these crucial industries.

Objectives:

- to evaluate and use the theory of queuing to optimize wait times in different service systems.
- to look at the ideal amount of servers needed to cut down on wait times without going over budget.
- to research the dynamics of lines in India's railway reservation and hospital bed allocation systems.
- to investigate the use of bulk queuing models in public utilities.

Methodology:

The M/M/S model, a multi-server queuing model, is used in this study to examine steady-state solutions and service dynamics in scenarios including large arrivals and services.

Literature

Literature Review:

Since its origin, queuing theory has seen tremendous development, mostly as a mathematical tool for waiting line analysis and optimization. The first works that contributed to queuing theory were written by A.K. Erlang in order to evaluate telephone call traffic, Erlang created the first queuing models in 1909 (Erlang, 1909). Erlang's contributions established the groundwork for subsequent research in queuing theory, which later spread to a number of industries, such as service management, healthcare, and telecommunications (Cooper, 1981). In today's complex systems, where service demand frequently outpaces available resources, queuing theory is crucial.

Applications of queuing theory in healthcare administration, particularly with reference to hospital bed allocation, have been well documented in recent years. When bed demand outpaces supply during emergencies, queuing algorithms have been demonstrated to maximize bed management (Green, 2006; Harper & Shahani, 2002). Harper and Shahani (2002) talk about how efficient queuing systems might improve patient outcomes by cutting down on wait times in intensive care units and emergency departments. The study emphasizes how hospitals may effectively distribute resources by using queuing theory to identify bottlenecks and forecast peak times. Furthermore, Green (2006) emphasizes the significance of queuing theory in cost-effective healthcare administration by pointing out that hospitals can decrease patient wait times by implementing queuing models without requiring expanding their physical infrastructure.

Another area where queuing theory has been widely used is in the transportation industry, specifically in railway systems. One of the biggest railway networks in the world, Indian Railways, has a difficult time controlling the massive number of people waiting in line to book tickets (Bhattacharya & Bandyopadhyay, 2015). The use of queuing theory to railway ticket reservation systems was studied by Bhattacharya and Bandyopadhyay (2015), who also suggested methods for increasing reservation efficiency through queue management. According to their findings, the ticketing process may be greatly streamlined by applying queuing theory, which will increase customer satisfaction and lower operating expenses.

The M/M/S model is a popular model in queuing theory that is applied to systems with Poisson arrival rates, exponential service times, and numerous servers (Gross & Harris, 2008). When examining complicated queuing systems, where customers are likely to leave the line if wait times are too lengthy, the M/M/S model is especially helpful (Bhat, 2015). According to Bhat (2015), the M/M/S model offers insights into how to best allocate server numbers to reduce wait times and service expenses. This is especially important in resource-constrained industries like public transit and healthcare.

The use of queuing theory has been extended in recent research to bulk queuing scenarios, when units or customers arrive in groups and necessitate certain service procedures. Queuing models aid in resource allocation based on bulk arrival patterns in bulk service systems, such as freight logistics and vehicle dispatch, guaranteeing more seamless and effective operations (Daganzo, 1997). Bulk queuing models are particularly useful in transportation systems, where traffic congestion can be controlled by modifying service prices to account for fluctuating demand levels, according to Daganzo (1997).

Results and discussion:

Because of their hectic lifestyles, people in today's culture frequently lack patience. No matter where you are—at a bus stop, a traffic light, a bank, a post office, a gas station, that elevator, or somewhere else—every moment you spend standing in line is problematic. In the course of our daily lives, we frequently encounter lines or lineups. Spending time waiting for something to happen is a waste of time. Cutting down on waiting time is one of our objectives. Generally speaking, in order to decrease waiting time, additional financial commitments are needed. To make an informed decision on whether or not to invest, one must comprehend how the investment will impact the amount of time spent waiting. Another name for a line is a well-organized group of individuals or cars waiting for their turn to be served or move forward. Queues enable institutions or organizations to adopt an organized approach in order to provide services in an effective manner. If managed effectively, the social phenomenon of the construction of a line may have positive effects on society. This will provide the greatest potential benefits for both those providing service and those waiting. One way to characterize a queuing system is the movement of units searching for services, joining or entering the queue when service is not immediately available, and exiting the system after receiving service. In

order to systematically analyze waiting lines or lineups, queuing theory examines them from a mathematical standpoint. The arrival in the line is one of several interconnected processes that can be thoroughly investigated according to the theory. Numerous performance metrics can be derived and calculated using the theory, including the average wait time in the queue or system, the expected number of people waiting for or receiving service, and the likelihood that the system is in specific states, such as full, empty, having a server that is available, or requiring a specific wait time. Because of its many significant applications, the theory has been well documented in the literature on probability theory, operations research, management services, and industrial engineering. Numerous applications have been thoroughly examined. The flow of customers is the primary concern of a wait situation in the setting of a service station. If the server is available, customers can be served right away when they arrive; if not, they might have to wait until the service is ready. These kinds of circumstances frequently arise throughout routine interactions. The commercial service system, the transportation service system, the industrial service system, and the social service system are only a few of the numerous instances of queuing systems. It makes sense to assume that customers will have to wait in order to receive service because there are only so many servers accessible. Adding more servers to the system will reduce the anticipated waiting time, but the cost of delivering the service will go up as a result. This makes it feasible to determine the ideal number of servers that will lower the total costs associated with waiting times and service. Although this component makes sense conceptually, it is challenging to precisely calculate the number of cases per unit of waiting time in operational settings. It is possible to conclude from this that queuing theory is not an optimization strategy. It serves as a tool that offers more detailed information on the topic under discussion. When a group of people go to a restaurant, they can get service together. This kind of circumstance is referred known as queuing. A Poisson stream of customers arrives in groups and is served at a counter in batches of varying sizes, according to the basic concept of bulk service. The Poisson technique is followed in doing this. Until the queue size reaches or exceeds a predetermined threshold (less than or equal to the capacity), the server remains idle. At that point, all of the customers are handled together. The server remains inactive till that time. There is significant value in implementing queuing applications across various public utilities. A thorough examination of system loss can be successfully carried out by applying queuing theory to telephone networks. The concepts derived from queuing theory have been effectively applied in a number of contexts, such as the processes followed by airplanes when

they reach their destination and the allocation of hospital beds. In particular, the goal is to conduct a thorough analysis of the bulk queue phenomena and the applications of queuing theory to issues related to hospital bed distribution and Indian train ticket reservations. An M/M/S system's steady state solutions will be produced, and waiting time characteristics will be investigated. The behavior of clients who are hesitant to use the system because of the limited area will be considered in these solutions. The outcome of a particular server scenario will be assessed as a unique case that might or might not be consistent with the results of earlier studies. The main focus of the study will be the vehicle dispatch mechanism, particularly when there are large numbers of arrivals and services. An examination of the queue dynamics associated with these possible scenarios will also be part of the study. When it comes to distributing beds across the several wards they have, hospitals face significant challenges. For emergency ambulances and intensive care units (ICUs), prompt service is crucial because any delay in their arrival could have disastrous consequences. In order to minimize the amount of time spent waiting for admission, the current study will apply queuing theory to determine the most efficient bed distribution among the various hospital wards. Either of the two options given can be used to reserve train tickets for Indian Railways. One method involves the individual purchasing an advance ticket in person at the counter established by Indian Railways, while the other option involves making a reservation online. Furthermore, we advise using automated teller machines (ATMs) to make cash withdrawals and to purchase train tickets.

Applications of the Clearly Defined Configuration Duration Model,

During the setup phase, we have given special attention to several things, such as cleaning the bucket filter, servicing the burner assembly, cleaning the pre-heating tank, servicing the oil pump, and servicing the nozzle. As part of the service we offer for the burner assembly, we will focus on maintaining the burner assembly's component parts. The components include the electrode and the viewing glass in addition to the divisor plate. The fact that these parts are being worked on while the heat pack is operating is being addressed. Ash accumulates on the surface of these components as a result of this. The ash has been fully extracted and thoroughly cleaned. Therefore, it is crucial that these components be ash-free prior to the treatment being administered. Cleaning the bucket filter is the next important step that needs to be finished in

order to eliminate the contaminants that are present in the furnace oil. There are a significant number of pollutants present since the furnace oil is in a mixed state. As a result, the filter undergoes maintenance and is cleared of any potential contaminants. The Nozzle, a comparatively small but crucial part that needs to be kept clear of contaminants right away, comes next. The servicing must be finished before the process. If this issue is not resolved, it could lead to back pressure, which could possibly result in a process flaw. The pre-heater tank is now being maintained by keeping an eye on the water level indicator that is located within the thermo pack. The water level begins to drop as the process progresses. Therefore, ensuring comprehensive service is crucial. The hot pumps undergo maintenance to ensure that there are no leaks throughout the system. The reduction of leaks is essential to ensuring a decrease in the quantity of raw materials utilized. Furthermore, a number of auxiliary services are being managed concurrently, including pump couplings and current utilization.

Following the completion of this procedure, the used oil will pass through the pipe and move on to the first stage, which involves heating the thermo pack. The oil undergoes a battery of tests to determine how much carbon it contains before being permitted to enter the thermo pack. The procedure is made possible when the oil contains only a trace amount of carbon, which is permitted to enter the thermo pack. If the amount of carbon in the oil is significantly higher than normal, the information was removed since it was irrelevant to the process being performed. Within the framework of queueing theory, this elimination process is referred to as "reneging." Only the oil with a lower carbon concentration is permitted to circulate within the thermo pack. The heating process starts when this threshold is crossed. A thermo pack serves a similar purpose as a gas stove, which is a tool we use on a daily basis. The furnace oil is changed into a liquid state throughout this procedure. It is permitted for the furnace oil to pass through the nozzle at a high velocity. The nozzle's output is rather small in relation to its size. The furnace oil has changed into a gaseous state by the time it exits the nozzle. The electrode is in charge of initiating the furnace oil's burning while also acting as a source of ignition. This then leads to the subsequent heating procedure. The oil that is contained inside the thermo pack begins to warm up. We believe that for oil to be successful, it must reach a temperature of 185 degrees Celsius. The temperature of the thermo pack (oil) varies with the changing of the seasons. We are able to ensure that the temperature of the surrounding air during the winter months stays constant because the thermo pack's boiler temperature can reach 210 degrees

Celsius. The thermo pack reaches a temperature of 185 degrees Celsius in the warmer months. The way a thermo pack works is comparable to how a gas stove works. When the fuel is furnace oil and the electrode is lighter during the heating process, the substance that needs to be heated is olive oil.

The second step of the process must be completed when certain services have been completed. Overall, the process is made more efficient by using these services. In this specific case, there are 72 molds in total. There are four different groups of the 72 molds, which are called presses. We have a total of 18 presses in our production plant. The top pump and the lower pump are the two pumps that comprise every single press. These pumps go through a thorough service procedure before proceeding to stage 2. Numerous pipe connectors are included with these pumps. Every press has undergone the appropriate maintenance. The lower pump has been adjusted to operate between 145 and 150 degrees Celsius, while the upper pump has been calibrated to operate between 150 and 155 degrees Celsius, reflecting our consumption patterns. Hot oils are sent to the presses designated for them via pipelines. For equipment, these pipelines are essential passageways. The pipes traverse the entire module. They serve a similar purpose to that of the veins that are located inside the human body. The thermo pack serves as the main component of this process in addition to being the pumping mechanism.

The system won't be operational on days with no working hours. As the oil travels through the pipes, it begins to transform from a liquid to an amalgam condition. Because of the oil, there is a barrier that prevents fresh oil from getting into the pipes. As a result, they operate as an obstruction inside the pipes to draw attention to themselves. Before employing toluene, the oil particle that has been there for a long time must be removed. The ability of this specific toluene to dissolve the oil ensures that there won't be any obstructions in the pipeline. This process guarantees that there are no dust particles present and continuously improves the oil flow. It is imperative that this preprocessing step be completed before proceeding to stage 2. Additionally, as the next step in the preparation process, we will surround the mold with a melamine box. Heat transmission to the exterior of the box will be successfully prevented by the mold's construction, which ensures that the warmth remains concentrated inside the mold. The temperature inside the melamine box reaches 150 degrees Celsius when it is introduced,

whereas the temperature outside stays at about 60 degrees Celsius. This prevents the worker from being exposed to temperatures over usual in the vicinity of the mold.

Only when the temperature surrounding the mold is satisfied is it permitted to apply the oil. The mold's temperature will progressively increase if oil is continuously flowing around it. The result of this evolution is the raw materials' inability to solidify. At this stage, we will employ a three-way valve to address the problem. This valve is composed of three separate components. The parts of the system are the input valve, the exit valve, and the drain valve. This specific valve has a diaphragm attached to it that is connected to the intake valve. Included are the header line's inline design and the intake valve, respectively. A link exists between the diaphragm system and the temperature sensor. The temperature surrounding the mold is measured using this temperature sensor, and the data is subsequently transmitted to the diaphragm system. The amount of oil that is permitted to surround the mold depends on the temperature. This vital oil is moved to the output valve with the aid of the input valve. The oil flow is halted and any remaining oil is let to move from the input valve to the drain valve after the mold's temperature is kept steady. A connection has been established between the header line's perimeter and the drain valve. This is the manner in which the renegeing begins. The three-way valve is the instrument used to keep an eye on the hot oil. A diaphragm that is positioned in the valve's crown position is one of the assembly's parts. Both the intake valve and this diaphragm are connected. The applied pressure, which is then connected to the compressor, can be used to observe the diaphragm's movement.

In the second stage of preparation, we facilitate the flow of oil throughout the mold. The oil circulation must be as consistent and dependable as feasible. An increase in the oil flow indicates that the mold's temperature has risen. The mold's temperature decreases while it is at a low temperature. This method proceeds as follows, in a manner similar to heating an object with a candle. The temperature of the candle increases when the object is placed inside it for a while. When the temperature of an object is continuously lowered, it does not increase and stays in the same state as when it was initially produced. In this specific context, the procedure is inverted. The mold's temperature is elevated during the process by the oil that surrounds it constantly circulating. Following a period of inactivity, the mold's temperature begins to drop. All things considered, a temperature of 150 degrees Celsius is kept around the mold in

accordance with the previously indicated specifications. The part of the product that is utilized for molding is the press, which also serves as the die in the system. The bottom die does not move, but the higher die can be adjusted to fit the circumstances. The die's location and movement will be assessed as part of the product evaluation. Each individual unit's construction takes a considerable amount of time and frequently involves heavy weights. The artwork is given to the media to help with a more effective assembly during the die replacement process.

In the initial step of the procedure, the oil's temperature is raised using a thermo pack. It was found that the oil's temperature would rise in direct proportion to the volume of the workpiece that needed to be molded. Temperature dissipation, fluid leaks, pump failures, and other related problems are some of the difficulties we face when working on the oil gearbox. The oil will be permitted to proceed to the next stage of the procedure, which is now in progress, after the barriers have undergone the required maintenance. The mold will be deemed to be in the mold once the temperature of the hot oil surrounding it equals the temperature of the mold. By properly regulating the temperature inside the mold, we ultimately succeed in achieving our objective. The temperature is raised in order to maximize the casting process's efficiency. Even now, a considerable number of industrial businesses continue to use this process.

Conclusion:

There is much promise for improving service efficiency and cutting wait times by implementing queuing theory in a variety of public utilities, including hospitals and train reservations. The results of the study show that a high level of service may be maintained while waiting expenses are reduced through optimal server allocation. Decision-makers can evaluate and modify resources to satisfy demand, enhance customer happiness, and preserve operational efficiency with the use of the M/M/S model. Thus, queuing theory becomes a useful framework for effectively and economically tackling the problems of service system management.

In summary, the body of research highlights how queuing theory may be applied to a variety of problems in a wide range of industries. Queuing models allow enterprises to make data-driven decisions that improve efficiency and save operating costs, from managing Indian Railways ticket reservations to allocating hospital beds as efficiently as possible. By investigating the applications of queuing theory in certain high-impact situations, this study

seeks to add to the corpus of knowledge by providing insights into workable strategies for cutting wait times and improving service delivery.

References

- Bhat, U. N. (2015). An Overview of Queueing Theory: Analysis and Modeling in Practice. Birkhäuser.
- Bandyopadhyay, D., and Bhattacharya, R. (2015). Queueing theory applied to the train reservation system. 6(6), 108–112. International Journal of Scientific & Engineering Research.
- Cooper, Robert B. (1981). Overview of Queueing Theory. Elsevier.
- C. F. Daganzo. (1997). foundations of traffic operations and transportation. The Pergamon Press.
- Erlang, A. K. (1909). Probability theory and telephone talks. Matematik's Nyt Tidsskrift.
- L. V. Green (2006). Healthcare queueing analysis. Reducing Healthcare Delivery Delays in Patient Flow (pg. 281-307). Springer.
- Harris, C. M., and Gross, D. (2008). The basics of queueing theory. Wiley & Sons, John.
- Shahani, A. K., and Harper, P. R. (2002). modeling for hospital bed capacity planning and management. Operational Research Society Journal, 53(1), 11–18.