

VIRTUALIZED STOCHASTIC MULTI QUEUING MODEL BASED ON VMAX/M/G-MIN/I-CACHE - P λ M FOR REDUCING CUSTOMER IDLE TIME EXPECTATION ON SINGLE SERVER UTILIZATION

A. ANDREW MICHAEL¹ AND M. THIAGARAJAN

ABSTRACT. Queuing is tremendous model for reducing workloads or execution of process in various functionalities using Markovian General Distribution (M/G/D). It expects to reduce the time of accomplishment arriving on ordered or unordered representation of workload in server utilization. Due to maximum arrival of customer request on single server at a time based on M/G/I, the server dumps to create the scheduling problem because of multi task. The large scale stochastic model (LSSM) doesn't concentrate on this problem. To resolve this problem we modulated the distributed process into virtualized stochastic model (Vmax/M/G-min/I-cache) on multi queue setup model (P λ M) in single server utilization called Virtualized stochastic multi queuing model (VSMQM). The Virtual Resource queuing System is divided into four parts by creating single to multimodal service performance indicators in a shared distribution mode. Sorting the request buffer cache is created based on customer request arrival rate and resource output. To improve the use of model complex virtual resource sharing, different service performance outputs can achieve a different parameter input transformation on queuing theory.

¹*corresponding author*

2010 *Mathematics Subject Classification.* 68M20, 60K25.

Key words and phrases. Max queuing, Probability matrix, distributed queuing, scheduling, server theory.

1. INTRODUCTION

Queuing Theory is a type of probability theory used to describe more specialized mathematical models, as they wait in line. The systems that we actually use to deploy different types of representations use an array model. You can also find the right balance of sample service cost and waiting volume [1]. Delivering a queue system customer service consists of a service and client (s) waiting on the server (s). In recent years, server utilization data centers have seen an increase in energy consumption because of more customer in queuing request.

To switch Vm based virtualization to extend on number of queuing requests that whether the number of server depends the sleep mode or idle remains they terminates to reduce the execution. There are many applications due to $M / G / 1$ as a prototype for real systems or web nodes where service on these systems is sometimes exponentially distributed and cannot be explained. The web server is set up in $M / G / 1$ model format. The measured speedup [2] is derived using the controlled Poisson distribution to describe the service time.

However, some of them can be ignored as model power efficiency methods that do not overhead the array system and the migration of the VM. We will analyze the performance of the system by modeling the energy efficient work planning mechanism by switching to a server-like / sleep mode in a multi-server vacation. First, we introduce a server / sleep mode transition and energy efficient work planning mechanism in virtualized data centers. Then queuing theory probabilities are task arrival markovian point executing remains to calculate the energy consumption and response efficiency. In this study, the network in sorting can be thought of as an arbitrary but systematic way in which the patient flows past the number of interconnected rows [3]. The sorting network can also be classified as closed open or mixed according to the different forms of the structure that the patient enters or leaves. In an open sorting network, patients are treated in one or more centers, entering the network from the outside and leaving the network. In comparison, in a closed queue network, no patients leave or enter networks.

2. PRELIMINARIES

Sorting Managing Tasks plays an important role in business process purpose. Analyst provides a powerful tool for evaluating the effectiveness of modeling and sorting methods in sorting. From the perspective of the number of tasks and the number of servers in the system sleep mode. The Poisson Perado Blasting Process (PPBP) has proposed a more realistic model of Internet traffic than its predecessor. In this model, bursts of data (e.g. files) are generated according to a Poisson procedure of parameter λ . These eruptions transmit any amount of Perato propagation, each of them a fixed rate R . According to the model, when people are active, new sources can start transmitting at any time, and we can rate any number of sources at the same time. This is a further generalization of the model when the burst length is normally distributed. In this case, the introduction of this model as a function of time fits in the evolution of the $M / G / \infty$ sorting system. An $M / G / \infty$ resource is at the same time active and equipped with busy m servers in the system. The sequence structure name $M / G / \infty$ is used as an alternate name by its very nature. Poisson processes can produce large number of independent events of sample events, each of which produces precisely low density spaced probability events. Such events can add up to several generations of phone and Internet traffic flows.

3. VIRTUALIZED STOCHASTIC MULTI QUEUING MODEL BASED ON $M/G/1$ - $P\lambda M$

Where the customer waiting time expects the service execution policy on multi state by splitting the single server using Service Speed Layer Distribution System Time after time $M1$ is replaced by an $M / G / 1$ array system be into visualization. Consider a $P\lambda M$ service policy. An explicit representation of the work flow constant distribution is presented in the transcendental argument. The $P\lambda M$ policy is an add on distribution process based a task of releasing water from a DAM which arrives and distribution based on time of increasing level to be releasing. It will change this cache representation of single queuing theory policy and introduce a new service policy for Max- $M/G/D$ -Min/ A -Cache sorting methods. The server was initially unresponsive. When the client arrives, server 1 starts virtualization process on parallel distribution based on customer request

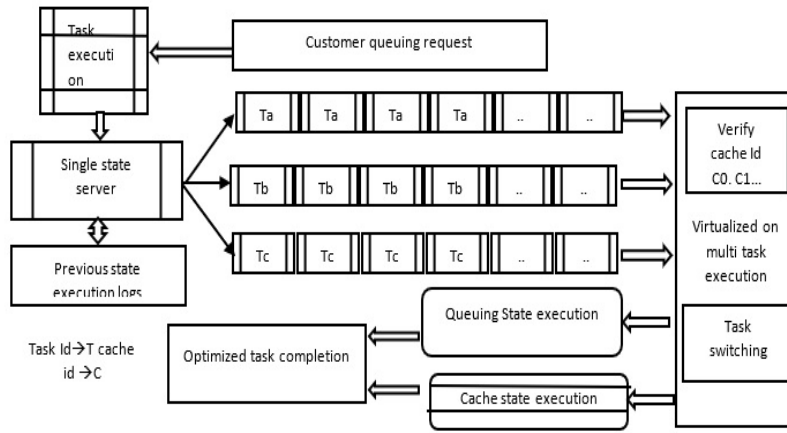


Figure 1: Process Virtualized stochastic multi queuing model based on M/G/1- $P\lambda M$

on multi queuing to improve the service speed and reduce the server workload. We modulated the distributed process into Virtualized stochastic model ($V_{max}/M/G\text{-min}/I\text{-cache}$) on multi queue setup model ($P\lambda M$) in single server utilization called Virtualized stochastic multi queuing model (VSMQM. When the customer speeds, according to the function G is distributed based on the poisson process independently represented as queuing theory on VM migration as T_s and the average M If the system's workload level exceeds $\lambda > 0$, the server instantly switches to server1- M into virtualized server and his speed is up to M . Restart the server to serve at 1 speed when another client arrives. The above Figure 1 shows the Process Virtualized stochastic multi queuing model based on M/G/1- $P\lambda M$. To deal with such a system, we have maintained a service speed of $M1$ when the system's workload level λ has been exceeded until the system is empty after the allocated time for the random system. Suppose that the speed of service you want is changing. If the server is free before the time is up, then he keeps his service up to speed 1. Setup time is also indicated by S , independently exponentially averaged $1 / \lambda$, assuming that the workload is distributed. The $\rho < m < 0$ system stability 1. We need to show the critical performance indicator in the sorting model, G-min/I-cache achieve our main objective is to govern the steady queue dissemination of the work flow in steady state queuing theory on Max state.

Queuing System: M stands for Markovian or memory less. The first M denotes arrivals following a Poisson process, the second M denotes service time following

exponential distribution. This refers to a single server and refers to infinite system capacity. Queuing System: This system is a type of queue with at most N customers allowed in the system. This system is a type of queue with at most N customers allowed in the system. The ultimate purpose of organizing a system review is so that it can be informed and clever decisions can be eliminated by management in order to understand the underlying process behavior. Applying the queuing concept is an attempt to reduce costs and reduce system delays due to disabilities. Various methods have been proposed to solve the problem of sorting. In this study, FIFO searches between two time-varying groups in a single server Markov sorting and unlimited order in a single server Markov sorting between service time discipline and unlimited client populations. As we have in the catering data, the estimated parameters for the same applies to this model. The sensitivity analysis system is evaluated for stability.

3.1. The M/G/1 queuing representation. The queuing Markovian is distributed, also depends on the time spent on the server for the client's departure, before further service. - This information is obtained from the variable N (D) which means that the faster service time, the service size already has no bearings (memory lessness), M (memory less): Bayesian arrival process, G (general) λ Severity: public load time distribution, average $S = 1 / \mu$ 1: single server, load $\rho = \lambda S$ depends ($S < 1$ in a fixed order), customizable in the system The number, $n(t)$, does not now constitute a Markov process. Unit transition probability $N = N - 1$ from state to state undergraduate = n, the average queue is not included. Server resembles the task arrival between the arrival and completion on target execution of V_m represented based on assign and terminate which can be found in presentations.

3.2. Customer Inter-arrival service on virtual queuing request. A single server channel offers Bayesian input, high-speed service, and a sorting system with unlimited capacity on a FCFS on basis of Max representation. We have taken from requests from various distorts be consider customer requirements (CI) like Arrival Time on clock (ATC), Intertribal Time IAT, Service response time (SRT).

3.3. Single server M/G/1 queuing on customer arrival. Let us consider a single server on V_{mm} extended queue $= \frac{\rho^2}{1-\rho}(V_m), MA/min/ - VM$ represented state equalized $Lq = \frac{\rho^2}{2*(1-\rho)}$ to generalize the level of arrival state by

CR	IAT	ATC	SRT		CRN	ATC	Req-Start	Duration	End
1	-	0	2		1	0	0	2	2
2	2	2	1		2	2	2	1	3
3	4	6	3		3	6	6	3	9
4	1	7	2		4	7	9	2	11
5	2	9	1		5	9	11	1	12
6	6	15	4		6	15	15	4	19

TABLE 1. Service time arrival and duration rate

its formulate queuing representation. $Lq = \frac{\rho^2(1+C_s^2)(C_a^2+\rho^2C_s^2)}{2*(1-\rho)(1+\rho^2C_s^2)} \rightarrow Ex(VM)$, representation by the arrival rate of task by mean rate vitruakl server λ and σ_a^2 inter arrival difference a is variance arrival rate: $VM \rightarrow C_a^2 = \frac{\sigma_a^2}{(1/\lambda)^2}$ Likewise, if μ execute service availability rate and σ_s^2 where s service time evaluation time, then: $Vm \rightarrow C_s^2 = \frac{\sigma_s^2}{(1/\mu)^2} * Vm$ the generalize the approximation as $Lq = \frac{\rho^2(C_a^2+C_s^2)}{2*(1-\rho)}$, where g is equalized max if state is max/ Min approximation $g = \exp\left(\frac{(1-\rho)(1-C_a^2)}{(C_a^2+4C_s^2)}\right)$ limit VM remains, when is $C_a^2 > 3$ Poisson distribution remains the intervals defining at the random state whether the customer or the task be arrived in that queue in arri9val and execution state of the queuing theory which depend of task arrival at a time.

Customer	Integrated Time [Minutes]	Appearance Time	Service Time [Minutes]	Starting time service	Time Service Ends	Queuing wait state [Minutes]	Customer spending time	Server idle hold state
1	-	0	4	0	4	0	4	0
2	1	1	2	4	6	3	5	0
3	1	2	5	6	11	4	9	0
4	6	8	4	11	15	3	7	0
5	3	11	1	15	16	4	5	0
6	7	18	5	18	23	0	5	2
100	5	415	20	416	418	1	3	0
Total	415		317			174	491	101

TABLE 2. Inter-arrival on service time distribution

This step monitors the project if the logical hyper dynamic scheduling request is non-volatile and updates each graph theory point to this cache time task switch P (process) and the program arrival rate d. The necessary steps to

deliver the expected vehicle speed can be performed at run time. So that a px and math program is written that will start executing it will perform read and write tasks at the time of execution. Since David Walsh is an active algorithm, the operation $td \rightarrow tt$ program is an intuitive number f.

Hence, by the inverse function $ts(x)$ and $x(ts) \rightarrow psd * trans$, Where the $Max \rightarrow ts(x) = ts(x) + \int_{x0}^x \frac{du}{v(u)} p(x \rightarrow vm)$ by similarly Delay request flow in that points vehicles arrival rate is estimated by continues chain rule.

Mean variance representation continues flow by max compared to other side of vehicles arrival rate, $Sx(v) = ts + \int_{v0}^v \frac{v}{v(v)} dv(VM)$ By the form of Fuzzy time logic dynamic scheduling; Max state precedence allotted to each Delay request represented four lanes. $td = \sum_{i=1}^m dx(i) * vm$ by the Initialization arrival rate $Av(ts)$ at the Max-state Transiting S with $tt = \sum_{j=1}^m dt(j)$ points to waiting arrival rate. $Lq = \frac{\rho^2(1+C_s^2)(C_a^2+\rho^2C_s^2)}{2*(1-\rho)(1+\rho^2C_s^2)}(VM)arravl$ at-cache($P \rightarrow id(VM)$), representation by the arrival rate of task by mean rate λ at virtualization and σ_a^2 inter arrival difference a is variance arrival rate: $C_a^2 = \frac{\sigma_a^2}{(1/\lambda)^2}$ Likewise, if μ execute service availability rate and σ_s^2 where s service time evaluation time, then: $C_s^2 = \frac{Vm \rightarrow \sigma_s^2}{(1/\mu)^2}$ the generalize the approximation as $Lq = \frac{\rho^2(C_a^2+C_s^2)}{2*(1-\rho)} + Cid$, where g is equalized max if state is max/ Min approximation detained from single level customer arrival on cache $g = \exp\left(\frac{-2*(1-\rho)(1-C_a^2)^2}{3p(C_a^2+C_s^2)}\right) * VM(Exec)$ repeat request Pid (Vid), when is $C_a^2 \leq 3$, and $g = \exp\left(\frac{(1-\rho)(1-C_a^2)}{(C_a^2+4C_s^2)}\right) * Vm$ as PID, when is $C_a^2 > 3$ whether queue cache has retained reduced time.

3.4. Queuing Performance evaluation. The number of customer arrival at that their utilization remains the single server process is high , the performance at constant state is idle weather task in queuing state .performance state is complete the task on reducing the queuing on Virtualization mode wheather its need at a time weather task are arrive at queue. Mean time estimation which the service as spend more time in order increasing utilization. Customers often combine quality service to wait longer. If you suffer from waiting too long, an option will be released once needed. Service limits extends the tasks depends the execution VM need only the consumption until remains to turn terminate. As execution on VM. The Percentage of usage measuring time server busy. Make sure management queue is subject to standby tolerance.Average waiting time

$= W = \sum_{i=0}^{i=n} \frac{\text{Waiting time in queue}}{\text{Number of customers}}$ To analyze the M/M/1 queue, average time on delay occurrence prior on task arrival and chose your choice Think of a customer. Since the "tag" is for this client, it goes in queue and looks in the opposite direction of late. The customer will travel through the queue and receive the service. Average total time is the average of all individual delay sources in the system.

3.5. Probability of Queuing arrivals. Number of intervention between the customer request that probability makes time process to wait p (wait) $= \frac{\text{Number of customer who wait}}{\text{Number of customers}}$ With N number of μ customers based on that time until t time executes at in arrival rate at $t=0$ on queuing system. The Poisson distribution is described as probability distribution where for a finite sorting system that cannot be extended to the right of infinity. This block is controlled by the state, $I = s$, the rest of this state's system, although it does not allow the system to enter this way. Proportion of server idle time $p(\text{idle server}) = \sum_{i=0}^{i=n} \frac{\text{Idle time of server}}{\text{Simulation run time}}$ The waiting line model is important for companies because they directly impact customer service awareness and service cost. Some functional areas have been affected by awaiting a decision on taxes. The wait-line system used by the treasurer is about cost.

3.6. Service arrival on execution. If the average usage of the system is low, it says that the waiting line design is inefficient and very expensive. Bad system design can lead to the acquisition of unnecessary capital to improve surplus personnel and customer service. Average service time $= S = \sum_{i=0}^{i=n} \frac{\text{Service time}}{\text{Number of customers}}$ Therefore, the probability that the array is occupied by an arrival event is simply the utility in terms of u , and the customer mean rate at waiting but not the arrival time is $Q - U$ receives, on average, a service time for each customer. Therefore, it is necessary to serve all customers who are waiting in line for the arrival time of the expected time is $(Q - U)s$. Average time between arrivals $= \lambda = \sum_{i=0}^{i=n} \frac{\text{Times between arrivals}}{\text{Number of arrivals}-1}$ be the consideration $E(\lambda) = \frac{a+b}{2}$ Average waiting time of those who wait $= \text{waited} = \sum_{i=0}^{i=n} \frac{\text{Waiting time in queue}}{\text{Number of customers that wait}}$ Customers need to have one or more taxes, and other service discipline issues, regardless of how many lines are in balance and how long to wait. Features of the Service. Simple queue system, each client offers only one server, how many servers. Customer mean rate spends in system $= t = \sum_{i=0}^{i=n} \frac{\text{Time customers spend in system}}{\text{Number of customers}}$ usually

the service system is considered random, exponentially distributed, and if there are multiple servers, it is generally considered that all servers are identical.

3.7. Changing phase on VM migration. This migrate the Single server into VM based on the task arrival the VM extents by needs that is dedicated to being part of the total service, rather than having to know the full service it offers rather than the server. Where they resembles the Queue represented by execution the Speed of service can be increased by spending extra resources. Generally the representation target of high limit is achieved at the Virtualized state. Let us consider the arrival time and the average spread of 10 minutes. There are two servers, each server has a service time of 20 minutes, a minimum of 10 minutes delay tolerance, a minimum wait time (i), an average number of queues (ii) and a uniform distribution. (v) Percentage average number, server is inactive. Results from different phenomenon simulations, known to be very accurate, show sequence average latency of 9.5693 minutes. M / M / G-VM approximation error.

$\lambda = \frac{1}{10}$, $C_s^2 = 1$ as defined as value, the service time be mean time $(10+20)/2=15$, follows $\mu = 115$, the service time distribution time, σ_a^2 will dependant $(20 - 10) 2/12 = 8.33$, $weather \rho = \frac{15}{2*10} = 0.75$, where $C_s^2 = \frac{\sigma_s^2}{(\frac{1}{\mu})^2} = \frac{8.33}{(15)^2} = 0.03$, by assigning the Queuing on representation the task arrival at the rate, $\rho_0 = 0.1453$, then $Lq = \frac{\rho^2(1+C_s^2)(C_a^2+\rho^2C_s^2)}{2*(1-\rho)^2} = \frac{0.1453(\frac{1}{15})^2*0.75}{2*(1-0.75)^2} = Wq = \frac{Lq}{\lambda} = 1.929 * 10 = 19.29$.

3.8. Continual Virtualization. Wq equivalent $\frac{(C_a^2+C_s^2)}{2}(19.29) * (1 + 0.03)/2 = 9.93$, by termination the target representation the arrival in queues on each task will be $Lq = Wq * \lambda = 9.93 * 1/10 = 0.993$, The average representation on Vm extended form is $\frac{|9.9376 - 9.5693|}{9.9376} * 100 = 3.07\%$, The mean estimation delay time is $W = Wq + \frac{1}{\mu} = 9.93 + 15 = 24.93$ min Number of system be virtualized by average mean time is $L = \lambda w = 2.2493$. virtual system that executes the task of evaluation to remains the execution state. From the above mean time delay reduces the waiting time in queue which manage the virtualized flow through the system faster. The Vm changes the priority rules based on task arrivals, max utilization become vulnerability of these tasks which wait a long time to change the priority. Single line queues model is more appropriate to change the VM to multi-line model and the customer is concerned about his reliability. It is a

single line guarantee that he does not anchor in an attempt to take advantage of another customer. This can easily add multiple queue models to a special server (quick execution). If changes are proposed, the impact of standby taxation on individual features is assessed by virtual machines. Changes in one area require changes in another the server reduces the waiting utilization.

4. RESULTS AND DISCUSSION

The Virtualized stochastic multi queuing model based on Vmax/M/G-min/I-cache - P λ M Single server is to predict between the request service transmissions depends the waiting, and a transmitting of workload for scheduling state. From the service disciplines, to compare the proposed system with Shortest Remaining Processing Time first (SRPT) as base to the Existing base and also associate Biggest-In-First-Served (BIFS). There is considerable correlation between the Max state of is carried through Time state intelligence in scheduling. Testing parameters consider the i3 process designed the program in visual frame work the arrival task are data packets which describes the specifics of the Queuing service values to estimate the demonstration of the planned method. This defines the Queuing service simulation parameters that are covered in max state scheduling. Average time to evaluate service in queue Average sacrificial arrival rate.

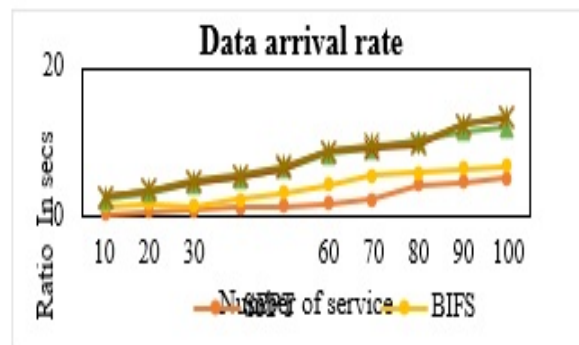


FIGURE 1. Analysis of Data rate

Figure 2 describes the clearly understood planned method. Consumes pointed out this of Data arrival rate.

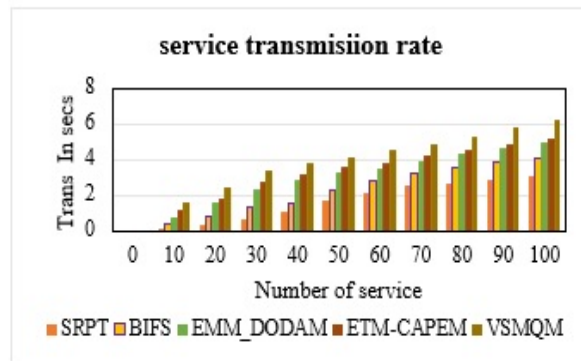


FIGURE 2. Analysis of Service transmission

Figure 3 describes the spread rate is calculated based on the number of services released at any given time. This is just a quick gauge of how much information you need to send broadcast technology with definitions. Higher propagation of proposed agricultural products than existing systems. The importance of queue service can be calculated by gradually adding a high-scheduling system across the chain to the multicast, enabling energy, information mystery and data processing to capture all information in the form of service arrivals. Thus, phase time logic can be used to optimize the resource life of the following locations. Figure 4 describes transmission Energy Performance Prepared by different

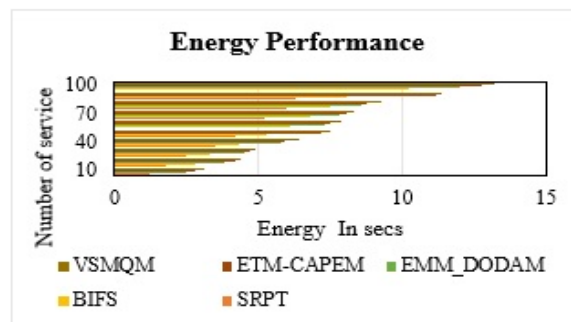


Figure 4 Energy Performance:

FIGURE 3

methodologies, and the performance of the proposed program increased the duration of energy. Higher efficiency than the proposed systems of energy levels

above existing systems the end of service output response in the proposed queue with increased duration.

5. CONCLUSION

From the results of the research to the fact that waiting lines can lead to many inefficiencies such as infinite growth, transient sorting systems (statistically analyzable). This means that their probability is higher than the assumption that it is for dissatisfied customers. The proposed Virtualized stochastic multi queuing model (VSMQM) improve the performance as well presenting the number of servers can be decreased to a main server extended auto Min 3 Vm extend and remains as to strengthen the effective order. Also, the speed with the current service has increased to $4.5660 \approx 5$ or more customers per minute. Our findings can be beneficial to improve the queuing theory performance as well delay rate, arrival performance, time consumption.

REFERENCES

- [1] J. CHEN, R. G. ADDIE, M. ZUKERMAN, T. D. NEAME: *Performance evaluation of a queue fed by a Poisson Lomax burst process*, IEEE communications letters, **19**(3) (2015), 367-370.
- [2] K. J. R. MARY, J. M. REMONA, R. RAJALAKSHMI: *Analysis of MX/M/1/MW V /BD queuing systems*, International Journal of Computer Applications, **141** (2016), 1-4.
- [3] T. SAKATA, S. KASAHARA: *Multi-server queue with job service time depending on a background process*, in Queueing Geory and Network Applications, 163-171, Springer, Berlin, Germany, 2016.

DEPARTMENT OF MATHEMATICS, LOYOLA COLLEGE, VETTAVALAM, THIRUVANNAMALAI-606754
 DEPARTMENT OF MATHEMATICS ST JOSEPH'S COLLEGE, Affiliated to Bharathidasan
 University, TRICHY-620002

DEPARTMENT OF MATHEMATICS ST JOSEPH'S COLLEGE, Affiliated to Bharathidasan University, TRICHY-620002