

The JMT Simulator for Performance Evaluation of Non-Product-Form Queueing Networks

M.Bertoli, G.Casale, G.Serazzi
Politecnico di Milano - DEI
Via Ponzio 34/5, I-20133, Milan, Italy
{bertoli, casale, serazzi}@elet.polimi.it

Abstract

This paper describes JSIM: the simulation module of the Java Modelling Tools (JMT), an open-source fully-portable Java suite for capacity planning studies. The simulator has been purposely developed to help both unexperienced and advanced users. Most of the difficult decisions that are needed in order to run simulations properly, such as the detection of the transient part of samples to be discarded, have been automated. The tool also provides guidance over the graphical design of the network and over the analysis and the plot of the results. What-if parametric analyses for parametric evaluation of complex systems are supported. Several features that increase the generality of the applications to capacity planning studies are provided, among them fork-join service centers, regions with finite capacity, state-dependent routing algorithms, priority classes and import of real workload distributions from log files.

1. Introduction

The availability of several simulation packages, either commercial or free, makes simulation one of the most commonly used techniques for performance evaluation of computer systems and networks. In general, evaluation techniques, that are methods by which performance evaluation indices are obtained, can be subdivided into two main categories: measurement (or empirical) techniques and modelling techniques. Empirical techniques require that the system or network to be evaluated exists and direct measurements of the evaluation target have to be taken. On the other hand, modelling techniques only require a model of the system. Modeling techniques are of two types: simulation and analytic. Among them, simulation is the most popular, as it applies to a wider variety of systems and does not require restrictive assumptions. However, in spite of their generality and ease of use, simulation models may

fail or produce non-accurate results. A first source of errors is related to the statistical techniques implemented in the simulator engine, e.g., the quality of the random number generator, the detection and removal of the transients, the algorithms used for confidence intervals and variance estimation. A second source of errors comes from users' mistakes, such as inadequate level of detail adopted to describe the target system, too short simulation time, errors in input parameter values and distributions, errors in output data interpretation and incorrect modeling of the characteristics of the target system. JSIM, the simulator described in the paper, aims at minimizing common mistakes by helping the average user in two ways. Firstly, critical statistical decisions, such as transient detection and removal, variance estimation, and simulation length control, have been completely automated, thus freeing the user from taking decisions about parameters s/he may not be familiar with. Secondly, a user-friendly graphical interface allows the user to describe, both the network layout and the input parameters. Furthermore, the graphical interface also provides support for advanced features like fork and join of customers, blocking mechanisms, regions with finite capacity constraints on population, state-dependent routing strategies, user-defined general distributions and import of log data. A module for what-if analyses with several types of control parameters, particularly useful in capacity planning, tuning and optimization studies is also provided. JSIM is fully implemented in Java and is available for download at the URL <http://jmt.sourceforge.net>. Since it is distributed under GNU GPL as an open source code, users can easily add new modules or integrate existing ones to customize the tool according to their needs.

The paper presents an overview of the simulation engine is given in Section 2. The supported statistical techniques are reviewed in Section 3. The salient JMT modelling features are described in Section 4. A case study focusing on the asymptotic behavior of queueing networks with finite capacity regions is described in Section 5. Conclusions are discussed in Section 6.

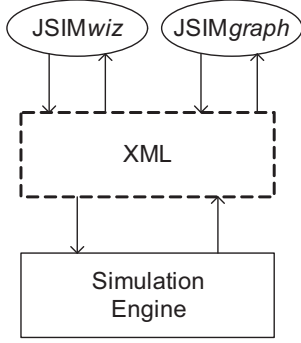


Figure 1. Simulator architecture.

2 Simulator Architecture

A graphical illustration of simulator architecture is given in Figure 1. The underlying design principle is to obtain a complete separation of the presentation and computational layers using XML. This approach has several benefits. For instance, it is possible to reuse the simulation engine within external applications; further, it simplifies the implementation of different graphical user interfaces. Concerning the last point, the JMT suite offers two simulation interfaces, called *JSIMwiz* and *JSIMgraph*, which are discussed later. In the rest of this section, we give an overview of the main functionalities of each architectural layer.

2.1 Discrete-Event Simulation Engine

The core module of the simulation engine is a discrete event calendar [2] that acts as a message broker, dispatching messages to simulation entities. Each significant event, e.g., the arrival of a new job to a queue or the departure of a job after service completion, is represented by a message with a specific identification code. When all current events have been processed, the simulation current time is moved forward to the next instant with an event in the calendar.

In the simulated network, each service center is composed of three entities, called *sections*, as shown in the top-most diagram of Figure 2. The engine demands to the user to specify through the graphical user interface the input, service and output sections to be considered in each service center involved in the simulation. The bottommost diagram shows the instantiation of the three sections in a queueing center. The queue *input section* is used in this example to receive incoming jobs. It implements the queueing buffer and the queueing discipline, e.g., first-come first-served (FCFS) or last-come first-served (LCFS), which selects the jobs to be processed by the *service section*. The service section simulates the service process, e.g., a single server with user-specified service time distribution. After service comple-

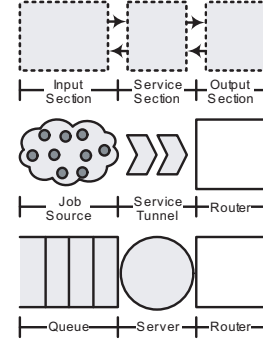


Figure 2. Service center sections.

tion, jobs are forwarded to a *output section* called *Router*, which sends them to the input section of another service center according to a user-specified routing strategy, e.g., Markovian routing.

As another example of application of this modular structure, we discuss the diagram in the middle of Figure 2, which represents a *source* service center. A source models the job arrival process for *open workloads*, i.e., flows of requests that arrive to the network according to a user-specified inter-arrival time distribution. The input section, called *job source*, is a pool of job objects that are sent, according to the selected distribution, to a fictitious server section, called *service tunnel*. This immediately forwards the job, with no delays, to the Router output section, which finally sends them into the network.

The main benefit of this modular representation of service centers is the ability of reusing the code of the same section in several service centers.

2.1.1 Simulation Coordination

The communication between entities is fundamental to coordinate simulations. In order to explain the message flows, we give an illustrative example. Let the service center *A* be a $M/M/1/\infty$ FCFS queue initially with three enqueued jobs. We denote the input, service and output sections of *A* respectively by *A.IN*, *A.SRV*, and *A.OUT*. Further, we call *immediate* a message that is processed by the simulator without increasing the simulation time. The example below is illustrated in the UML sequence diagram in Figure 3.

Message Flow Example

1. Initially, *A.IN* has been informed by *A.SRV* that the server is idle. Then, *A.IN* selects the next job to be serviced according to the implemented queueing strategy, and sends an immediate notification to *A.SRV*. The job moves to *A.SRV*;

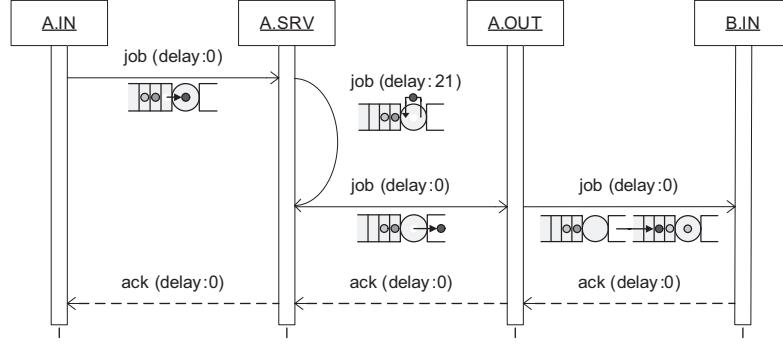


Figure 3. UML sequence diagram of a message flow example.

2. *A.SRV* receives the message and determines the job service time, e.g., $S_A = 21$, according to the user-specified service time distribution. Then, it sends to itself a notification with delay S_A ;
3. After S_A simulation time units, *A.SRV* receives the message and sends a new message to *A.OUT* to notify the service completion. The job moves to *A.OUT*;
4. *A.OUT* receives the messages, selects a routing destination, e.g., queue *B*, and notifies it to *B.IN*;
5. *B.IN* can either accept the new job or refuse it if the input buffer is full. The latter case is useful for models with blocking or finite capacity regions. If the job is accepted, *B.IN* sends an immediate acknowledgement to *A.OUT*. The job is now in *B.IN*;
6. *A.OUT* propagates the acknowledgement to *A.SRV*;
7. *A.SRV* informs *A.IN* with the acknowledgement that it is returned to the initial idle state.

Our messaging paradigm, which maintains a complete separation between sections, allows external contributors to develop new sections without the knowledge of other internal implementation details. In other words, in order to specify a new section it is sufficient to implement the correct messaging behavior.

2.1.2 Performance indices

The JMT simulator allows the computation of several performance indices: *Queue Length*, *Queue Time*, *Residence Time*, *Response Time*, *Utilization*, *Throughput*, *Drop Rate*, *System Response Time*, *System Throughput*, *Number of jobs* and *System Drop Rate*. Each performance index can be estimated for a particular job class or as an aggregated measure over all classes.

In order to collect samples, we use a special data structure referred to as *JobInfoList*. This is associated to each entity of interest, i.e., a section, a service center, or the entire network. The *JobInfoList* logs the arrival time of jobs and then, for each index of interest, feeds a statistical analyzer with the collected data. Each statistical analyzer estimates the requested index using spectral analysis and transient removal methods described in Section 3.

2.1.3 Control of Simulation Experiments

The user has several options for controlling simulation experiments. By default, the simulation is automatically stopped when at least one of the following criteria is satisfied by all performance indexes:

- both confidence interval estimates satisfy user requirements. This is implemented in JSIM through the concept of *maximum relative error*, which is equivalent to “the relative precision of the confidence interval” in [14], and stands for maximum acceptable ratio ϵ between the half-width of the confidence interval estimate (for the requested significance level α) and the estimated mean. For instance, setting $\epsilon = 0.1$ and $\alpha = 0.05$ imposes a simulation stop when the half-width of the estimated 95% confidence interval is no more than 10% of the non-transient sample mean.
- the number of collected samples exceeds a user-specified maximum ;
- the elapsed time exceeds a user-specified threshold;
- user aborts computation.

However, in some cases the user may be interested in performing *long-run simulations*. This feature is useful, e.g., for models with heavy-tail distributions, where large delays occur with non-negligible probability, and thus may not be observed if the simulation stops too early. For these cases,

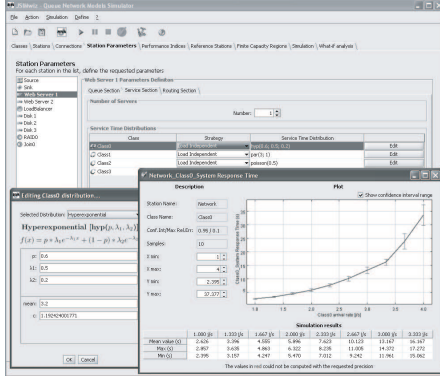


Figure 4. JSIMwiz screenshot.

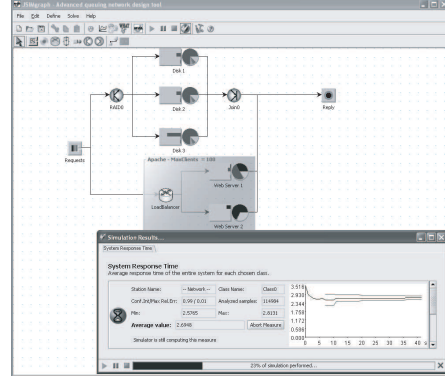


Figure 5. JSIMgraph screenshot.

JSIM gives the possibility of disabling the automatic stop feature.

Another important feature of the JMT simulator is the support for parametric analyses using *what-if simulations*. In practice, the analyst chooses a control parameter among the number of jobs for closed classes, the percentage of jobs belonging to a certain class, the service times for a service center, or the seed of the random number generator. In this case, JSIM performs a user-specified number of simulations varying the control parameter, and finally plots the results with confidence intervals over the different experiments. We show an example of this simulation feature in the final case study. A screenshot of a what-if plot is given in the right dialog in Figure 4.

2.2 XML data layer and Graphical User Interface

The simulator uses an XML data layer to invoke simulation engine, and to obtain simulation current estimates and final results. Furthermore, the XML format is used to save user-specified models: the output file format describes not only service center configurations and network topology, but, if required, also graphical positions of entities in the graphical interface and each Java class used to implement the different sections. This allows the introduction of new features using the Java Reflection API without modifying the simulation engine code.

The two graphical user interfaces to the simulator engine, called JSIMwiz and JSIMgraph, are built on the top of the XML layer. JSIMwiz offers a simple wizard-based interface that guides through model parametrization. JSIMgraph, instead, gives an easy-to-use graphical layout that enables to draw the network using, e.g., drag-and-drop of predefined service centers. Screenshots of the two interfaces are given in Figure 4 and Figure 5.

3 Statistical Analysis of Simulation Results

Simulation results are analyzed using transient detection and confidence interval estimation algorithms. These techniques are executed online, and can call for a simulation stop if all accuracy requirements are met. Confidence intervals are computed online using spectral methods [9]. However, their effectiveness depends on the stationarity of the sample distribution. In fact, unless long-run simulations are considered, transient effects can significantly affect simulation results. Therefore, discriminating if a group of samples describes transient or steady-state system performance is important to maximize accuracy, since transient data has to be discarded. We implemented transient detection using the R5 heuristic [6] and the MSER-5 stationarity rule [17]. The resulting transient detection and removal flowchart is shown in Figure 6. We now briefly review these methods and the most significant implementation decisions.

3.1 R5 Heuristic

The R5 heuristic [6] detects the initial transient period by checking if the time series of samples crosses its mean value more than k times, where k is a user-specified parameter. We implemented this rule using the indications of Pawlikowski [14]. In order to limit time and space requirements, we periodically check transient termination using a period of 500 samples. If the new samples do not satisfy the crossing condition, 500 new samples are collected before a new check is performed. After matching the stopping criteria, the R5 heuristic is complete and the simulator performs a second transient detection using the MSER-5 rule. This is done because we noted that applying several rules empirically reduces the number of wrong detections.

Concerning the choice of the k parameter, Gafarian et al. [7] recommend $k = 25$ for $M/M/1/\infty$ models, while Wilson and Prisker [19] recommend $k = 7$ for $M/M/1/15$

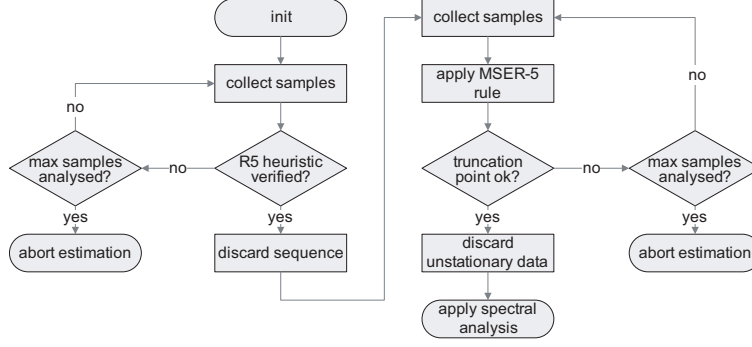


Figure 6. Transient detection and removal flowchart.

models. We inspected more than 50 single workload simulation models with randomly generated networks including up to 4 load-independent and load-dependent queues, and checked the behavior of utilization, throughput, queue-length and response time metrics. We observed that the values proposed in the literature did not correctly identified transients on some of the instances. In particular, there exist cases in which $k = 7$ gives an early detection during an unfinished initial transient ramp. Conversely, with the $k = 25$ value, the transient of the utilization metric tends to be very slow, and in some cases even hundreds of thousands of samples are discarded before moving to the MSER-5 rule. We found empirically that the $k = 19$ setting produces good results also on problematic instances.

3.2 MSER-5 Rule

MSER-5 rule is an identification method for the best truncation-point in a data sequence, discussed in [17]. Initially, we used Schruben’s test for this purpose [16], but the results were not always satisfactory, as observed also in [10]. Since the MSER-5 rule is not generally meant for online analyses, but instead works on a fixed data set, we adopted the approach of Robinson [15] which consists in getting online new samples unless the truncation point is detected before the half of the analyzed sample set. Our implementation uses circular lists of 5000 batches (25000 samples), and has constant access and computation times. We observed that using structures with increased size had no benefit on accuracy, but imposed additional computational overheads. In the simulator, the MSER-5 rule is reapplied periodically until the detection of the optimal truncation-point. We set the period equal to the number of samples discarded by the R5 heuristic.

3.3 Spectral analysis

The spectral analysis of Heidelberger and Welch [9] is a stable and computationally efficient method for computing simulation confidence intervals using variable batch sizes and a fixed amount of memory. In the JMT simulator, the method is run on the non-transient part of the sequence of samples. The spectral analysis is periodically run until the required confidence intervals are found. Nonetheless, it happens quite often that the portion of data immediately available after the MSER-5 rule is already sufficient to compute the required confidence intervals.

4 Non-Product-Form Modelling Features

In this section we give an overview of the supported modelling features of the simulation engine. We give a description of significant design choices for each feature.

4.1 Arrival and Service Processes

Nowadays performance models require heavy-tail distributions for arrival and service processes which can be modelled using, e.g., Pareto distributions. Exact results for models representing these effects are seldom available, and typically simulation is required. The JMT simulator supports several distributions for arrival and service processes. These include, among others, Pareto, Gamma, Hyperexponential and Erlang. Random number generation is based on the Mersenne Twister engine [13]. We also implemented a *LogReplayer* distribution that allows to import the random values from an external text file. Load-dependent service processes of arbitrary type can also be specified. These are useful to model devices with service times depending on the current queue-length. Moreover, load-dependence is required for hierarchical modelling and parametric analyses [5].

4.2 Fork-Join Service Centers

Fork-join service centers are employed to represent resources that can serve jobs in parallel [8], and are frequently used in storage, parallel and grid system modelling. Fork-join service centers are composed by $P > 1$ queues in parallel. Each time a job arrives to a fork-join service centers, it is split by a *fork* node into P sibling *tasks*. Each of them is assigned to a different parallel queue. After receiving service, jobs synchronize and merge at a *join* node before leaving the service center. Figure 7(a) shows an open network with a fork-join service center. Despite investigated for a long time, approximate solution techniques for networks with fork-join service centers are limited by simplifying assumptions, e.g., exponentiality of service processes.

In the JMT simulator, the user needs only to connect fork and join nodes to any subnetwork. The fork node is implemented as a service center with arbitrarily-chosen input section, with a *tunnel* service section, and with a special output section that sends tasks into all outbound connections. Each task is associated to a data structure which identifies the original job arrived at the fork node, and the number P of sibling tasks. The *join* node has a special input section which waits for all P tasks before forwarding to a tunnel section the merged job. The output section of the join is defined by the user.

4.3 Finite Capacity Regions

Models of simultaneous resource possession due to memory or software constraints often require *finite capacity regions* (see, e.g., [12]). These are subnetworks where the number of circulating jobs is constrained. *Shared* constraints impose an upper bound on the allowed number of jobs in the region regardless of their service class. *Dedicated* constraints, instead, limit the number of cycling jobs for a specific class. Jobs arriving to a full region enqueue in a *waiting buffer* outside the region. An illustration is given in Figure 7(b). The presence of the waiting buffer makes it difficult to obtain an analytical solution to models with finite capacity regions. Therefore, only approximation techniques have been developed [12]. However, simulation remains the most important analysis technique in presence of realistic workloads with multiple classes.

Finite capacity regions are implemented as follows. The waiting buffer is a service center with infinite capacity queue and tunnel service section. The output section implements the access control policy according to the user-specified shared and dedicated constraints. Waiting jobs are selected to enter the region according to a FCFS discipline. We point out that region access control is centralized, i.e., all arrivals are routed to the same waiting buffer. Currently, the simulator does not support multiple waiting buffers and

nested or intersecting regions. We also remark that, when a region is full, the user can force the simulator to drop arriving jobs.

4.4 Priority Classes

Priority modeling is required in numerous applications including, e.g., models of packet flows that are differentiated according to quality classes or in the analysis of scheduling rules. Priority models have studied since the early years of performance modelling, and several analytical results of interest are available [8].

In the JMT simulator, we support priority disciplines without preemption, namely priority FCFS and priority LCFS. The integration of preemptive methods is currently left as future work.

4.5 State-Dependent Routing Strategies

A key assumption of product-form queueing networks is that the behavior of a service center depends only on its current state, and not on the current state of the other queues. Sophisticated generalizations for integrating features depending on network state have been presented in the literature (e.g., [18]), but they are limited to specific classes of models. Thus, in general, one needs simulation to evaluate the impact of state-dependent routing techniques. JMT simulator supports the following state-dependent routing strategies: *Shortest Queue Length*, *Shortest Response Time*, *Least Utilization* and *Fastest Service*.

5 Case Study

We present a case study that illustrates the potentialities of the JMT simulator. We compare the asymptotic behavior of product-form network versus non-product-form networks with finite capacity regions. While the distribution of jobs for asymptotically large populations is well understood in product-form theory (e.g., [1]), little is known on non-product-form networks. Our contribution is to show some interesting differences in the asymptotic behavior of multiclass finite capacity networks.

Network Description We consider an example network composed by three queues and by a delay. The model has the topology shown in Figure 8. The Web application runs on a server modelled as a finite capacity region with two identical queues CPU1 and CPU2. The region accounts for the maximum number of simultaneous user sessions. We assume that the Web application concludes its activity by placing an order ($Class - 1$), or by calling a service ($Class - 2$), from a Backend

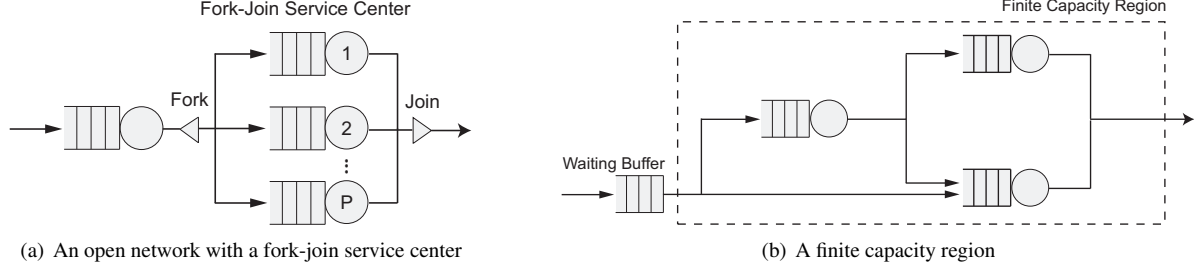


Figure 7. Some non product-form features supported by the JMT simulator.

server. Assuming negligible communication overheads, the completed requests return immediately to the clients. The average service time of requests at the different service centers is $S_{CPU1,Class-1} = 1.0$, $S_{CPU1,Class-2} = 10.0$, $S_{CPU2,Class-1} = 1.0$, $S_{CPU2,Class-2} = 10.0$, $S_{Backend,Class-1} = 10.0$, $S_{Backend,Class-2} = 2.0$, $S_{Delay,Class-1} = 10.0$, $S_{Delay,Class-2} = 25.0$. The routing probability from the delay station to each of the CPU queues is 0.5.

In our experiment, the population sizes N_1 and N_2 of the two workloads are changed while keeping fixed their sum to the constant value $N_1 + N_2 = 1000$ job. Denoting by $\beta_1 = N_1/(N_1 + N_2)$ and $\beta_2 = N_2/(N_1 + N_2)$, such that $\beta_1 + \beta_2 = 1$, the percentage of jobs belonging to the two classes, our study consists in a parametric analysis of the network behavior as a function of β_2 . Note that assigning a specific value to β_2 immediately implies $\beta_1 = 1 - \beta_2$. We evaluated the network for $\beta_2 = 0.0, 0.1, 0.2, \dots, 1.0$ collecting 10^5 samples for each server utilization in the following four scenarios:

- *Product-Form Case.* The finite capacity region is disabled, and the number of jobs that can enter CPU1 and CPU2 has no constraints. In order to meet product-form assumptions, we assume that all queues have the processor-sharing scheduling rule. The results reported for this case are exact values.

- *Shared Constraint Case.* In this case we assume that the aggregated number of jobs in CPU1 and CPU2 is less than or equal to 150, regardless of the class of membership. All queues have now the FCFS scheduling rule that usually considered in models with finite capacity regions.
- *Class-1 Dedicated Constraint Case.* This scenario is analogous to the Shared Constraint Case, but the 150-jobs bound is a limit on the class 1 population only.
- *Class-2 Dedicated Constraint Case.* This case is similar to the previous one, but the constraint is now placed on the class-2 population only.

Discussion of Experimental Results. The total mean utilizations of the CPUs and of the Backend servers in the different scenarios are given in Figure 9. It is well-known from the asymptotic theory of [1] that there exists a continuous interval of values for the population mix β_2 where both the CPUs and the Backend centers must reach maximum utilization $U_{CPU1} = U_{CPU2} = U_{Backend} = 100\%$. This interval, called *common saturation sector*, is highlighted in Figure 9(a)-(e), and is slightly smaller than the theoretical ones obtainable with the formulas in [1], which hold for asymptotically large populations, since we had to consider a *finite* total population¹.

By comparison with the other scenarios, we clearly see that the predictions of multiclass product-form models may be significantly different from those of models where we account also for the finite population constraints. In particular, we found that the shared and class-1 dedicated constraint case *do not exhibit common saturation sectors*. Thus, the CPUs and the Backend center simple interchange, for a certain value of β_2 , the role of “bottleneck” center limiting network performance.

The scenario with the class-2 dedicated constraint in Figure 9(d)-(h) has a behavior more similar to the product-form case. at the main difference of this model with respect to the

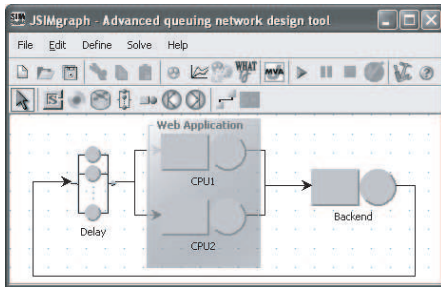


Figure 8. Case study: network topology.

¹The boundaries of the asymptotic common saturation sector can be determined with the JABA tool of the JMT suite [3].

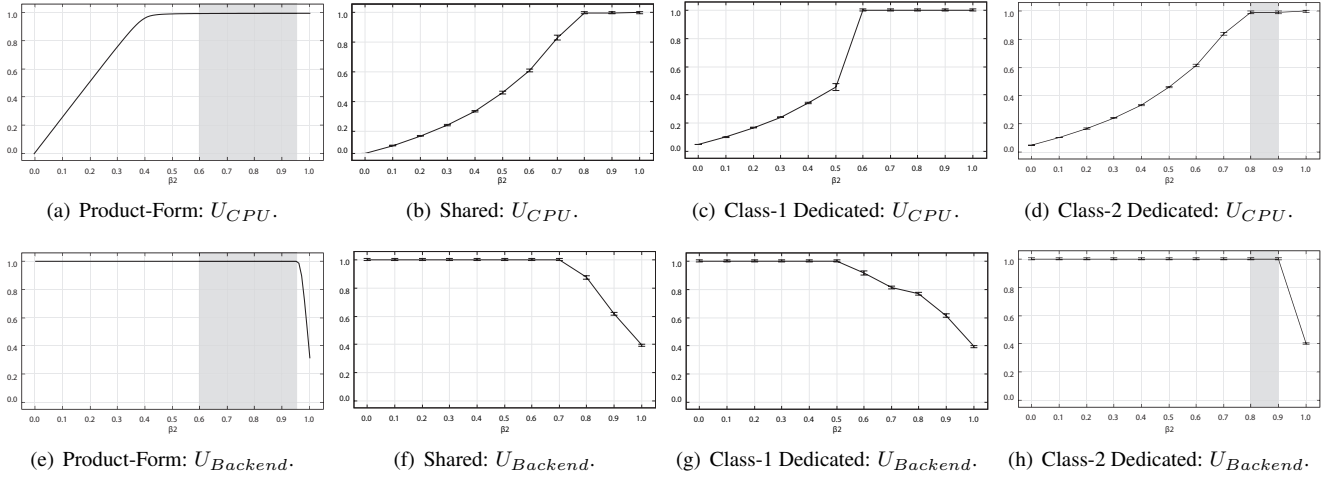


Figure 9. Server utilizations as a function of the population mix β_2 ($U_{CPU} = U_{CPU1} = U_{CPU2}$)

others is that, for $\beta_2 \in [0.8, 1.0]$, the class-2 jobs saturate both the CPUs and the Backend server. The only other case where there is a large number of jobs inside the CPUs is in the shared constraint case, but in this scenario we have that most of class-2 jobs reside in the waiting buffer of the finite capacity region.

In conclusion, the case study outlines the importance of the detail level used to model the target system. Hence, simulations should be carried out to determine which classes of non-product-form networks admit saturation sectors that can be approximated well under product-form assumptions.

6 Conclusions

The simulator can be downloaded, with the other tools included in the JMT suite [3], from the project homepage at the URL <http://jmt.sourceforge.net/>. Since the tool is released as a free open source software, the participation to the project of performance analysts and simulation researchers is welcome.

References

- [1] G. Balbo and G. Serazzi. Asymptotic analysis of multiclass closed queueing networks: Multiple bottlenecks. *Perf. Eval.*, 30(3):115–152, 1997.
- [2] J. Banks and J. Carson. *Discrete-Event System Simulation*. Prentice-Hall, Inc., New Jersey, USA, 1984.
- [3] M. Bertoli, G. Casale, and G. Serazzi. Java modelling tools: an open source suite for queueing network modelling and workload analysis. In *Proc. QEST'06*, 119–120, IEEE Press.
- [4] M. Bertoli, G. Casale, and G. Serazzi. An Overview of the JMT Queueing Network Simulator. Politecnico di Milano - DEI, TR 2007.2, 2007.
- [5] K. M. Chandy, U. Herzog, and L. Woo. Parametric analysis of queueing networks. *IBM J. Res. Dev.*, 19(1):36–42, 1975.
- [6] G. S. Fishman. Statistical analysis for queueing simulations. *Man. Sci., Series A (Theory)*, 20, 3:363–369, 1973.
- [7] A. V. Gafarian, C. J. Ancker, and T. Morisaku. Evaluation of commonly used rules for detecting “steady state” in computer simulation. *Nav. Res. Log. Q.*, vol. 25, 511–530, 1978.
- [8] G. Bolch, S. Greiner, H. de Meer, and K. Trivedi. *Queueing Networks and Markov Chains*. John Wiley and Sons, 1998.
- [9] P. Heidelberger and P. D. Welch. A spectral method for confidence interval generation and run length control in simulations. *Commun. ACM*, 24(4):233–245, 1981.
- [10] J. K. Preston White, M. J. Cobb, and S. C. Spratt. A comparison of five steady-state truncation heuristics for simulation. In *Proc. of the 32nd Winter Simulation Conference, SCS, Orlando, Florida, 2000.*, 755–760, 2000.
- [11] A. E. Krzesinski and P. Teunissen. Multiclass queueing networks with population constrained subnetworks. In *SIGMETRICS*, 128–139, 1985.
- [12] E. Lazowska, J. Zahorjan, G. Graham, and K. Sevcik. *Quantitative System Performance*. Prentice-Hall, 1984.
- [13] M. Matsumoto and T. Nishimura. Mersenne twister: a 623-dimensionally equidistributed uniform pseudo-random number generator. *ACM TOMACS*, 8(1):3–30, 1998.
- [14] K. Pawlikowski. Steady-state simulation of queueing processes: A survey of problems and solutions. *ACM Computing Surveys*, 22(2):123–168, June 1990.
- [15] S. Robinson. Automated analysis of simulation output data. In *Proc. of WSC'05 conf.*, 763–770., 2005.
- [16] L. Schruben. Detecting initialization bias in simulation output. *Operations Research*, 30:569–590, 1982.
- [17] S. C. Spratt. Heuristics for the startup problem, 1998.
- [18] D. Towsley. Queueing networks models with state-dependent routing. *JACM*, 27(2):323–337, 1980.
- [19] J. R. Wilson and A. A. B. Pritsker. A survey of research on the simulation startup problem. *Simulation*, 31, 2:55–58, 1978.