

Comparación de las prestaciones de latencia de los algoritmos de planificación DTable y DRR

Raúl Martínez, Francisco J. Alfaro, José L. Sánchez

Resumen—Una componente clave para redes con soporte de calidad de servicio (Quality of Service, QoS) es el algoritmo de planificación que se emplea en los enlaces de salida. Un planificador ideal implementado en una red de altas prestaciones con soporte para QoS debería satisfacer dos propiedades principales: buena latencia y simplicidad. El algoritmo Deficit Round Robin (DRR) es conocido por tener una complejidad computacional y de implementación muy baja. Sin embargo, dependiendo de la situación sus prestaciones de latencia pueden ser muy malas.

Por otro lado, los algoritmos de planificación basados en tabla tratan de proporcionar unas prestaciones de latencia aceptables con una complejidad baja. Algunas de las últimas propuestas de tecnologías de red como Advanced Switching e InfiniBand definen en sus especificaciones uno de estos planificadores. Sin embargo, los planificadores que proponen no funcionan adecuadamente con paquetes de tamaño variable y tienen el problema de ligar la asignación de ancho de banda y latencia.

Nosotros hemos propuesto un nuevo planificador basado en tabla, al que hemos llamado Deficit Table (DTable), que funciona adecuadamente con paquetes de tamaño variable. Además, hemos propuesto una metodología para configurar este planificador que desdota la ligadura entre ancho de banda y latencia.

Hemos evaluado las prestaciones del planificador DTable y la metodología de configuración. Los resultados de las simulaciones muestran, como esperábamos, que podemos asignar el ancho de banda a los distintos flujos con una cierta independencia de los requisitos de latencia que se les esté tratando de proporcionar. Sin embargo, no hemos comparado las prestaciones de latencia proporcionadas por el planificador DTable con las de ningún otro algoritmo de planificación. En este artículo, comparamos las prestaciones de este planificador con las prestaciones del algoritmo DRR.

Palabras clave—Evaluación de prestaciones, algoritmos de planificación, calidad de servicio, requisitos de ancho de banda y latencia.

I. INTRODUCTION

Current packet networks are required to carry not only traffic of applications such as e-mail or file transfer, which does not require pre-specified service guarantees, but also traffic of other applications that require different performance guarantees, like real-time video or telephony [10]. The best-effort service model, though suitable for the first type of applications, is not so for applications of the other type [11]. Even in the same application, different kinds of traffic (e.g. I/O requests, coherence control messages, synchronization and communication messages, etc.) can be considered, and it would be very interesting that they were treated according to their pri-

ority. Therefore, high performance packet networks need to enable Quality of Service (QoS) provisioning. The provision of QoS in computing and communication environments is currently the focus of much discussion and research in industry and academia. A key component for networks with QoS support is the output scheduling algorithm, which selects the next packet to be sent and determines when it should be transmitted, on the basis of some expected performance metrics.

An ideal scheduling algorithm implemented in a high performance network with QoS support should satisfy two main properties: good delay and simplicity. The design of a traffic scheduling algorithm involves an inevitable trade-off among these properties. Many scheduling algorithms have been proposed. Among them, the “sorted-priority” family of algorithms are known to offer very good delay [14]. However, their computational complexity is very high, making their implementation in high-speed networks rather difficult. In order to avoid the complexity of the sorted-priority approach, the Deficit Round Robin (DRR) algorithm [12] has been proposed.

The DRR algorithm associates each flow¹ with a *quantum* and a *deficit counter*. The quantum assigned to a flow is proportional to the bandwidth assigned to that flow. The sum of all the quanta is called the frame length. The deficit counter is set to 0 at the beginning. The scheduler visits sequentially each flow. For each flow, the scheduler transmits as many packets as the quantum allows. When a packet is transmitted, the quantum is reduced by the packet size. The unused quantum is saved in the deficit counter, representing the amount of quantum that the scheduler owes the flow. At the next round, the scheduler will add the previously saved quantum to the current quantum. When the flow has no packets to transmit, the quantum is discarded, since the flow has wasted its opportunity to transmit packets. The main advantage of the DRR scheduler is its computational simplicity. Recent research in the DiffServ area [3] proposes the DRR as a feasible solution for implementing the Expedited Forwarding Per-hop Behavior [5]. However, the main problem of this algorithm is that its delay depends on the frame length. Depending on the situation, the frame can be very long, and thus, the latency would be very bad.

On the other hand, in the table-based schedulers instead of serving packets of a flow in a single visit

Departamento de Sistemas Informáticos. Universidad de Castilla-La Mancha. 02071 - Albacete (España). {raulmm, falfaro, jsanchez}@dsi.uclm.es

Este trabajo ha sido parcialmente financiado por la CICYT con el proyecto TIC2003-08154-C06-02, por la Junta de Comunidades de Castilla-La Mancha con el proyecto PBC-05-005-1 y por el Ministerio de Educación y Ciencia con una beca FPU.

¹In this paper we will use the term *flow* to refer both to a single flow or to an aggregated of several flows with similar characteristics.

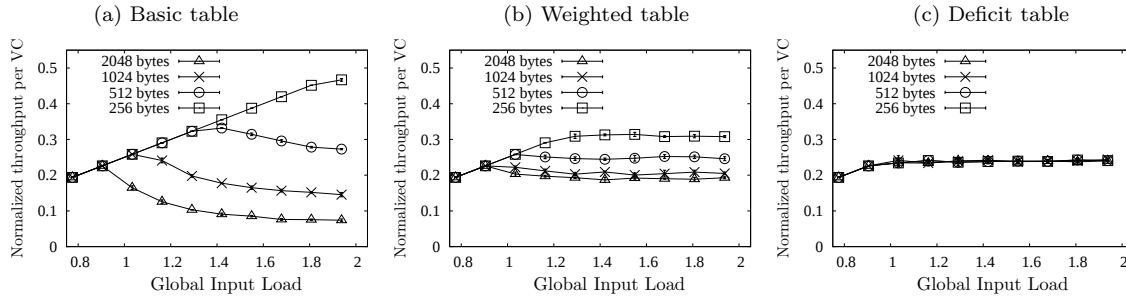


Fig. 1. Performance of several table-based schedulers for flows with different packet size.

per frame, the service is distributed throughout the entire frame. This approach is followed by two of the last high performance network interconnection proposals: Advanced Switching (AS) [1] and InfiniBand [6]. These table-based schedulers can provide a good latency performance with a low computational complexity [4], [9], [2]. However, these schedulers do not work properly with variable packet sizes and face the problem of bounding the bandwidth and latency assignments [9], [2]. In [8] we reviewed these problems and proposed a new table-based scheduler, that works properly with variable packet sizes. This scheduler defines an arbitration table in which each table entry has associated a flow identifier and an entry weight. The methodology that we propose to configure this scheduler set the maximum distance between any consecutive pair of entries assigned to a flow depending on its latency requirement. Moreover, it assigns the entry weights depending on the bandwidth requirements. This methodology permits us to decouple partially the bounding between the bandwidth and latency assignments.

In [8] we evaluated the performance of the DTable scheduler and its decoupling configuration methodology. Simulation results showed, as expected, that the weight assigned to the different VCs fixes the bandwidth they receive, independently of the number of table entries or the maximum distance between them, while the latency performance comes from the separation between any consecutive pair of table entries assigned to the VC. However, we did not compare the latency performance obtained with the DTable scheduler with the latency provided by any other scheduling algorithm. In this paper, we compare the latency performance of the DTable scheduler with the latency performance of the DRR scheduler.

The structure of the paper is as follows: In Section II, we review the DTable scheduler and our methodology to decouple the bandwidth and latency assignments. Details on the experimental platform and the performance evaluation are presented in Section III. Finally, some conclusions are given.

II. THE DEFICIT TABLE SCHEDULER

The main problem of the table-based schedulers proposed until now is that they do not work in a proper way with variable packet sizes. If the average packet size of the different flows is different, the bandwidth that the flows obtain may not be proportional to the number of table entries. Figure 1

shows the performance of various table-based schedulers when there are four Virtual Channels (VCs) in the network. Note that we use VCs to aggregate flows with similar characteristics and make the arbitration at a VC level, as it is the case in AS and InfiniBand. The four VCs have the same number of assigned table entries (the same bandwidth reservation). Moreover, we inject an increasing amount of traffic at the same rate in all the VCs. However, the traffic injected in each VC has a different packet size. The simulated architecture is the same as that used for the performance evaluation in Section III.

Figure 1(a) shows the case of a basic table scheduler similar to the AS table scheduler, which is cycled through and when a table entry is selected, a packet from the VC indicated in that entry is transmitted regardless of the packet size. As can be observed, when using the basic table scheduler, the VCs obtain a very different bandwidth because the traffic that traverses each VC has a different packet size. Therefore, although the same number of packets from each flow will be transmitted, the amount of information will not be the same.

The InfiniBand's arbitration table works in a similar way than the AS table. However, it adds a weight to each entry. This weight indicates the amount of information to be transmitted from the VC associated to the table entry each time that the entry is selected. This weighted table solves the problem only partially because it allows a packet to be transmitted that requires even more weight than the remainder of a given table entry (exhausting them). Figure 1(b) shows the performance of a weighted table that works in this way. We have assigned all the entries the same weight: 2176 bytes (34 units of 64 bytes). As can be seen, it presents a better performance than the basic table scheduler, but not an optimum performance.

In [8] we proposed a new table-based scheduling algorithm that works properly with variable packet sizes. (as can be seen in Figure 1(c)). We called this algorithm Deficit Table scheduler, or just DTable scheduler, because it is a mix between the already proposed table-based schedulers and the DRR algorithm. The scheduler works in a similar way than the DRR algorithm but instead of serving packets of a flow in a single visit per frame, the service is distributed throughout the entire frame. The computational complexity of this new scheduling algorithm is slightly higher than the complexity of the DRR scheduler.

The DTable scheduler defines an arbitration ta-

ble in which each table entry has associated a flow identifier and an *entry weight*. Moreover, each flow has assigned a *deficit counter* that is set to 0 at the beginning. When scheduling is needed, the table is cycled through sequentially until an entry assigned to an active flow is found. A flow is considered active when it stores at least one packet and the link-level flow control, if exists, allows that flow to transmit packets. When a table entry is selected, the *accumulated weight* is computed. The accumulated weight is equal to the sum of the deficit counter for the selected flow and the current entry weight. The scheduler transmits as many packets from the active flow as the accumulated weight allows. When a packet is transmitted, the accumulated weight is reduced by the packet size.

The next active table entry is selected if the flow becomes inactive or the accumulated weight becomes smaller than the size of the packet at the head of the queue. In the first case, the remaining accumulated weight is discarded and the deficit counter is set to zero. In the second case, the unused accumulated weight is saved in the deficit counter, representing the amount of weight that the scheduler owes the queue. The weights are usually expressed in flow control credits.

We set the minimum value that a table entry can have associated to the Maximum Transfer Unit (MTU) of the network. This is the smallest value that ensures that there will never be necessary to cycle through the entire table several times in order to gather enough weight for the transmission of a single packet. Note that this consideration is also made in the DRR algorithm definition [12].

In [2], we explained how to configure table-based schedulers (in that case for InfiniBand) to provide bandwidth and latency guarantees. In order to provide a flow with a minimum bandwidth, the number of table entries assigned to that flow must accomplish with the proportion of desired egress link bandwidth. In order to provide a flow with maximum latency requirements, the maximum separation between two consecutive table entries devoted to that flow must be fixed. By fixing this separation, it is possible to control the maximum latency of a network element crossing. This is because this distance determines the maximum time that a packet at the head of a flow queue is going to wait until being transmitted. Therefore, given a maximum number of hops, we can control the maximum end-to-end latency.

This way of assigning the entries of the table faces the problem of bounding the bandwidth and latency assignments. If a maximum separation between any consecutive pair of table entries of a flow is set, a certain number of them are being assigned, and hence a minimum bandwidth, to the flow in question. In that way, to assign to the most latency-restrictive flows a small amount of bandwidth is not possible because lower distances must be used for them. This can be a problem because the most latency-restrictive traffic does not usually present a high bandwidth require-

ment. Moreover, we cannot assign more bandwidth to the flows than that provided for the number of table entries.

In [8] we have also proposed a methodology to configure the DTable scheduler to decouple, at least partially, the bounding between the bandwidth and latency assignments. With this methodology we set the maximum distance between any consecutive pair of entries assigned to a flow depending on its latency requirement. Moreover, we can assign the flows with a bandwidth varying between a minimum and a maximum value that depends not only on the number of table entries assigned, but also on two table configuration parameters. We have called these parameters w and k .

Supposing an arbitration table with N entries in a network with a certain MTU, the w parameter determines the maximum weight M that can be assigned to a single table entry in function of the MTU: $M = MTU \times w$. The k parameter determines the total weight that can be distributed between all the table entries. We call this value the *bandwidth pool*: $pool = N \times MTU \times k$. The total number of weight units (a weight unit is equivalent to a flow control credit) from the bandwidth pool that the table entries of a flow have assigned fixes the bandwidth that the flow has actually assigned.

The w and k parameters fix the minimum bandwidth $min\phi_i$ and the maximum bandwidth $max\phi_i$ that can be assigned to the i^{th} flow depending on the number of table entries n_i that it has assigned:

$$min\phi_i = \frac{n_i \times MTU}{pool} = \frac{n_i \times MTU}{N \times MTU \times k} = \frac{n_i}{N} \times \frac{1}{k}$$

$$max\phi_i = \frac{n_i \times M}{pool} = \frac{n_i \times MTU \times w}{N \times MTU \times k} = \frac{n_i}{N} \times \frac{w}{k}$$

Summing up, the DTable scheduler is a table-based scheduler that is able to deal properly with variable packet sizes and considers the possibility of a link-level flow control mechanism. Moreover, with our configuration methodology we can provide a flow with latency and bandwidth requirements in a partially independent way.

III. PERFORMANCE EVALUATION

In [8] we evaluated the performance of the DTable scheduler and its decoupling configuration methodology. Simulation results showed, as expected, that the weight assigned to the different VCs fixes the bandwidth they receive, independently of the number of table entries or the maximum distance between them, while the latency performance comes from the separation between any consecutive pair of table entries assigned to the VC. However, we did not compare the latency performance obtained with the DTable scheduler with the provided by any other scheduling algorithm. In this section, we compare the performance of the DTable scheduler with the performance of the DRR scheduler. For this purpose, we have developed a detailed simulator that allows us to model the network at the register transfer level,

TABLA I
TABLE CONFIGURATION. $N = 64$, $k = 2$, $w = 4$.

VC	#entries	%entries	$\min\phi_i$	$\max\phi_i$
D2	32	50	25	100
D4	16	25	12.5	50
D8	8	12.5	6.25	25
D16	4	6.25	3.125	12.5
D32	2	3.125	1.5625	6.25
D64	1	1.5625	0.78125	3.125
D64'	1	1.5625	0.78125	3.125
Total	64	100	50	200

following the AS specification. Note, however, that our proposals can be applied to any interconnection network technology.

A. Simulated architecture

We have used a perfect-shuffle Bidirectional Multi-stage Interconnection Network (BMIN) with 64 endpoints connected using 48 8-port switches (3 stages of 16 switches). In AS any topology is possible, but we have used a MIN because it is a common solution for interconnection in current high-performance environments. The switch model uses a combined input-output buffer architecture with a crossbar to connect the buffers. Virtual output queuing has been implemented to solve the head-of-line blocking problem at switch level.

In our tests, the link bandwidth is 2.5 Gb/s but, with the AS 8b/10b encoding scheme, the maximum effective bandwidth for data traffic is only 2 Gb/s. We are assuming some internal speed-up (x1.5) for the crossbar, as is usually the case in most commercial switches. AS gives us the freedom to use any algorithm to schedule the crossbar, so we have implemented a round-robin scheduler. The time that a packet header takes to cross the switch without any load is 145 ns, which is based on the unloaded cut-through latency of the AS StarGen's *Merlin* switch [13].

A credit-based flow control protocol ensures that packets are only transmitted when there is enough buffer space at the other end to store them, making sure that no packets are dropped when congestion appears. AS uses Virtual Channels (VCs) to aggregate flows with similar characteristics and the flow control and the arbitration is made at VC level. The MTU of an AS packet is 2176 bytes, but we are going to use 2048 bytes (a power of two) for simplicity but without losing generality. The credit-based flow control unit is 64 bytes, and thus, the MTU corresponds to 32 credits.

The buffer capacity is 32768 bytes ($16 \times \text{MTU}$) per VC at the network interfaces and 16384 bytes ($8 \times \text{MTU}$) per VC both at the input and at the output ports of the switches. If an application tries to inject a packet into the network interface but the appropriate buffer is full, we suppose that the packet is stored in a queue of pending packets at the application layer. Regarding the latency statistics, a packet is considered injected when it is stored in the network interface.

TABLA II
BANDWIDTH CONFIGURATION.

VC	ϕ_i	DTable		DRR
		E. w.	T. w.	Quantum
D2	25	32	1024	256
D4	25	64	1024	256
D8	25	128	1024	256
D16	12.5	128	512	128
D32	6.25	128	256	64
D64	3.125	128	128	32
D64'	3.125	128	128	32
Total	100		4096	1024

B. Simulated scenario and scheduler configuration

We have defined 7 VCs with different distances between consecutive entries in the arbitration table. In a real case we would assign the traffic flows to these VCs depending on their latency requirements. Note that we are going to consider the requirements of a VC as the requirements of the traffic that is going to be transmitted using that VC. We have called these VCs D2, D4, D8, D16, D32, D64, and D64', indicating the distance between any pair of consecutive table entries. Therefore, D2 has more strict latency requirements than D4, D4 than D8, and so on. A table of 64 entries has been used in the simulations.

In order to allow the decoupling between the latency requirements of the VCs and the bandwidth assigned to them, we have used our methodology, assigning to the k parameter a value of 2 (the bandwidth pool is 2 times the MTU multiplied by the number of entries), and the w parameter a value of 4 (each table entry can be assigned a maximum weight of 4 times the MTU). Table I shows the number of entries assigned to each VC, the percentage of entries that this entails, and the minimum and maximum bandwidth that can be assigned to each VC using the decoupling methodology.

Table II shows the actual bandwidth configuration of the DTable and DRR schedulers. This table shows the amount of bandwidth ϕ_i that we have assigned to each VC. Regarding the DTable scheduler, it shows the total weight (T. w.) that we have distributed among the table entries of each VC and the weight assigned to each table entry (E. w.) of each VC. For example, in order to assign 25% of bandwidth to the D2 VC, 1024 of the 4096 credits of the bandwidth pool must be assigned to it. Therefore, 32 credits have been assigned to each one of its 32 table entries. Regarding the DRR scheduler, it shows the quantum assigned to each VC. The D64 and D64' VCs, which are the VCs with the minimum bandwidth requirement, are assigned a quantum that corresponds to 32 credits (the MTU), which ensures that at least one packet is going to be transmitted when a VC is selected. The rest of VCs are assigned a proportional quantum.

Note that the DRR frame length for this scenario is 1024 flow control credits. In the DTable scheduler that length is 4096 flow control credits. However, in the DTable case, the quantum assigned to each VC is distributed all along the frame length.

We are going to inject an increasing amount of

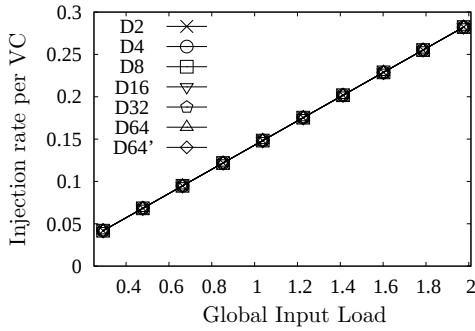


Fig. 2. Normalized injection rate per VC.

traffic of all the VCs and study the throughput and latency performance of both schedulers at different network load levels. The traffic load is composed of self-similar point-to-point flows of 1 Mb/s. The destination pattern is uniform in order to fully load the network. The packets' size is governed by a Pareto distribution, as recommended in [7]. In this way, many small-sized packets are generated, with an occasional packet of large size. The minimum payload size is 56 bytes, the maximum 2040 bytes, and the average 176 bytes, which represents enough packet size variability. The AS packet header size is 8 bytes. The periods between packets are modelled with a Poisson distribution. Figure 2 shows the normalized injection rate of the aggregated of flows associated with each VC.

C. Simulation results

Figure 3 shows the average values and the confidence intervals at 90% confidence level of ten different simulations performed at a given input load. For each simulation we obtain the normalized average throughput, the average message injection latency, and the maximum message injection latency of each flow. No statistics on packet loss are given because, as has been said, we assume a credit-based flow control mechanism to avoid dropping packets. We obtain statistics per VC aggregating the throughput of all the flows of the same VC, obtaining the average value of the average latency, and the maximum latency of all the flows. Note that the maximum latency shows the behavior of the flow with the worst performance.

Regarding the throughput performance, Figure 3 shows the normalized throughput results per VC of both schedulers. Note that both schedulers provide a similar throughput performance. As we can see, when the load is low, all the VCs obtain the bandwidth they inject. However, when the load is high (around 90%) the VCs do not yield a corresponding result, obtaining a bandwidth proportional to their assigned bandwidth. Note that the VCs do not obtain all the bandwidth that they were supposed to have assigned because the network is not able to provide 100% throughput.

Regarding the latency performance, Figure 3 shows that when the load is very low, all the VCs present a similar low latency. This is because at this load level there are few packets being transmitted

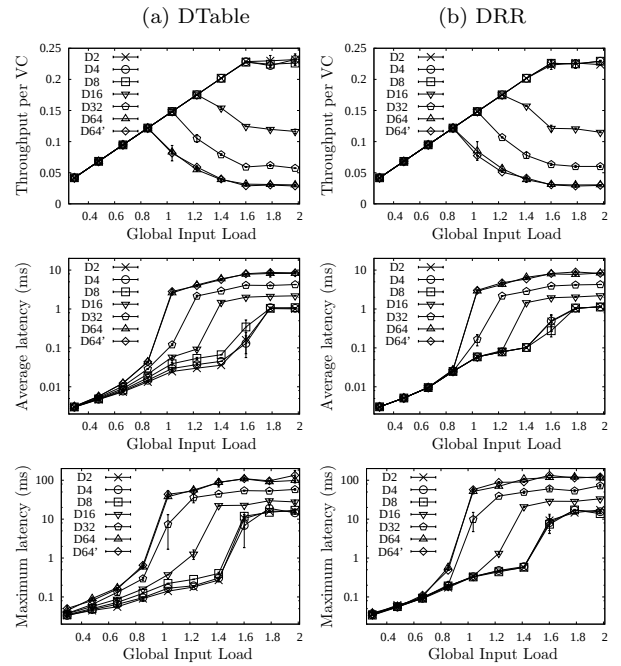


Fig. 3. Normalized throughput per VC of the different scheduling scenarios.

through the network, and thus, there are few conflicts between them. However, when the load increases, the latency also increases because some packets must wait in the buffers until others have been transmitted. It is at this point when the scheduling algorithm assumes an important role and the VCs obtain a different latency depending on the scheduling algorithm. However, when the load of the VC begins to outstrip its throughput, the latency starts to grow very fast. This is because the buffers used for that VC begin to be full. Finally, the buffers become completely full and the latency stabilizes at a given value which depends on the buffers' size and the bandwidth assigned to that VC, but not on the scheduling algorithm.

Note that when using the DRR algorithm, all the VCs obtain a similar latency performance until a VC reaches the point when its load begins to outstrip its throughput. In that point, the latency of that VC grows very fast and obtains a different latency performance. This happens for all the VCs as load grows. However, when using the DTable scheduler, all the VCs, including those with the same bandwidth assignment, obtain a different latency performance depending on the separation between any consecutive pair of their table entries. The smaller the distance, the better latency performance they obtain.

These different latency performance behaviors are explained by the fact that the maximum time that a packet at the head of a VC queue is going to wait until being transmitted is different depending on the scheduler algorithm. In the case of the DTable scheduler, we can control this time by controlling the maximum separation between any consecutive pair of entries assigned to the same VC. In the case of the DRR algorithm, the latency performance depends more on the frame length than on the quantum that each VC has been assigned. This is because when the quan-

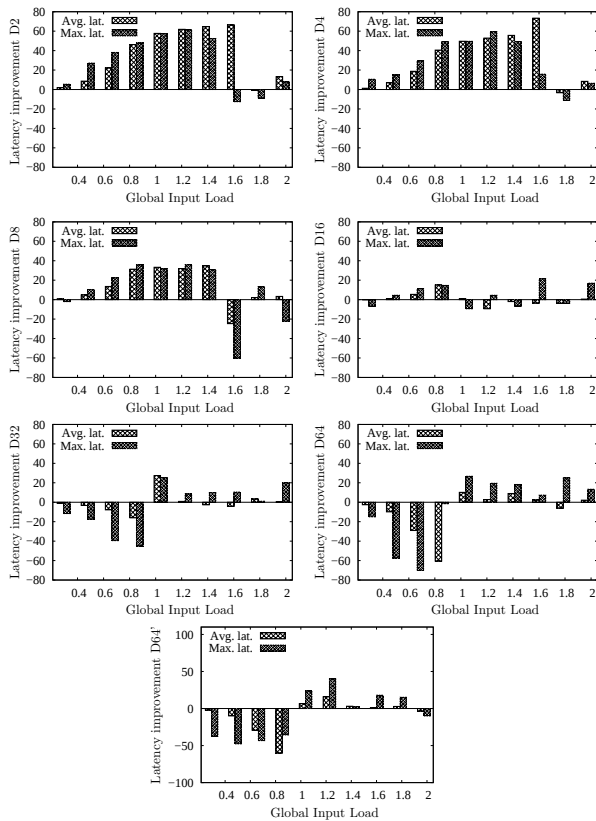


Fig. 4. Average and maximum latency improvement of the DTable scheduler over the DRR scheduler.

tum for a VC has been expended sending packets, all the frame must be cycled through before sending more packets of the same VC.

Finally, Figure 4 shows the percentage of improvement on the average and maximum latency provided by the DTable scheduler over the provided by the DRR algorithm. Note that, as stated before, when the injection of a VC is higher than its throughput, the buffers are almost always full. At this moment, the bandwidth that each VC is assigned outweighs in importance the distance of the entries in the arbitration table. This is the reason why the figures show that the percentage of improvement for those points becomes zero or does not stabilize in a clear value. Therefore, the relevant load level interval in order to compare both schedulers is from the point where the load starts to be high until the point where the VC saturates. Analyzing this load interval we can see that the DTable scheduler provides a quite better average and maximum latency for the D2, D4, and D8 VCs than the DRR scheduler. It provides a slightly better latency for the D16 VC and a worse latency performance for the D32, D64, and D64' VCs.

Summing up, the DTable and the DRR scheduling algorithms provide the same throughput performance, but a different latency performance. The DTable scheduler, which has a computational complexity slightly higher than the DRR algorithm, provides the most preferential VCs (those which have been assigned a shorter distance between any consecutive pair of entries) with a better latency performance than the DRR algorithm. However, it provides the least preferential VCs with a worse latency than the DRR algorithm.

IV. CONCLUSIONS

The main problem of the Deficit Round Robin (DRR) algorithm, which is known to have a very simple computational complexity, is that its delay depends on the frame length. Therefore, if the frame is very long the latency would be very bad. Moreover, we cannot differentiate the flows taking into account their latency requirements because all the flows obtain a similar latency performance.

We have proposed a new scheduling algorithm, the DTable scheduler, that solves in a high degree these DRR problems with only a slightly higher computational complexity than the DRR algorithm. With the DTable scheduler we can define several categories of traffic, with not only different bandwidth requirements but also different latency requirements.

In this paper we have compared by simulation the performance of the DTable and DRR schedulers. Simulation results show that the DTable scheduler provides the most preferential flows with a better latency performance than the DRR algorithm. However, it provides the least preferential VCs with a worse latency than the DRR algorithm. Therefore, with the DTable scheduler, we can provide a different level of latency performance to the VCs, prioritizing those VCs with higher latency requirements. This is not possible with the DRR algorithm.

REFERENCIAS

- [1] Advanced Switching Interconnect Special Interest Group. *Advanced Switching core architecture specification. Revision 1.0*, December 2003.
- [2] F. J. Alfaro, J. L. Sánchez, and J. Duato. QoS in InfiniBand subnetworks. *IEEE Transactions on Parallel and Distributed Systems*, 15(9):810–823, September 2004.
- [3] S. Blake, D. Back, M. Carlson, E. Davies, Z. Wang, and W. Weiss. An Architecture for Differentiated Services. Internet Request for Comment RFC 2475, Internet Engineering Task Force, December 1998.
- [4] H. M. Chaskar and U. Madhow. Fair scheduling with tunable latency: A round-robin approach. *IEEE/ACM Transactions on Networking*, 11(4):592–601, 2003.
- [5] A. Charny et al. Supplemental information for the new definition of EF PHB (expedited forwarding per-hop-behavior). RFC 3247, March 2002.
- [6] InfiniBand Trade Association. *InfiniBand architecture specification volume 1. Release 1.0*, October 2000.
- [7] R. Jain. *The art of computer system performance analysis: Techniques for experimental design, measurement, simulation and modeling*. John Wiley and Sons, Inc., 1991.
- [8] R. Martínez, F.J. Alfaro, and J.L. Sánchez. Decoupling the bandwidth and latency bounding for table-based schedulers. *International Conference on Parallel Processing (ICPP)*, August 2006.
- [9] R. Martínez, F.J. Alfaro, and J.L. Sánchez. Providing Quality of Service over Advanced Switching. *International Conference on Parallel and Distributed Systems (ICPADS)*, July 2006.
- [10] P. L. Montessoro and D. Pierattoni. Advanced research issues for tomorrow's multimedia networks. In *International Symposium on Information Technology (ITCC)*, 2001.
- [11] K. I Park. *QoS in Packet Networks*. Springer, 2005.
- [12] M. Shreedhar and G. Varghese. Efficient fair queueing using deficit round robin. In *SIGCOMM*, pages 231–242, 1995.
- [13] StarGen. *StarGen's Merlin switch*, 2004. http://www.stargen.com/products/merlin_switch.shtml.
- [14] D. Stiliadis and A. Varma. Latency-rate servers: a general model for analysis of traffic scheduling algorithms. *IEEE/ACM Transactions on Networking*, 1998.