

High Performance Dynamic Voltage/Frequency Scaling Algorithm for Real-time Dynamic Load Management[☆]

Javier O. Coronel*, José Simó Ten**

*University Institute of Control Systems and Industrial Computing (AI2)
Polytechnic University of Valencia
Camino de Vera s/n, 46022, Valencia, Spain*

Abstract

Modern cyber-physical systems assume a complex and dynamic interaction between the real world and the computing system in real-time. One of the challenges of future systems is to address the separation of modal control design from the deployment of executable code. In such systems it would be difficult, if not impossible, to analyze off-line all the possible combinations of processor loads. For this reason, it is worthwhile attempting to define new flexible architectures that enable computing systems to adapt to potential changes in the environment.

Separating design from code deployment could have a significant impact on the overall performance and real-time constraints of a system. Task migration approaches may contribute to on-line redistribution and reallocation of the workload – depending on each situation. In this context, we will consider a dynamic environment where an embedded and networked system operates with support for task migration and processor frequency scaling. We will assume that the system knows where and when tasks are allocated. To guarantee that the temporal requirements of the system are accomplished, we will perform a feasibility analysis when tasks are changed. A new processor speed (frequency scaling) is also computed to reduce energy consumption.

We present a novel algorithm to analyse task admission and processor frequency assignment. Additionally, we perform several simulations to evaluate the comparative performance of the proposed approach. This evaluation is made in terms of energy consumption, task rejection ratios, and real computing costs. The results of simulations show that from the cost, execution predictability, and task acceptance points of view, the proposed algorithm mostly outperforms other constant voltage scaling algorithms.

Keywords:

dynamic voltage scaling, task migration, real-time scheduling, power consumption, feasibility analysis.

1. Introduction

Modern cyber-physical systems (CPS) assume a complex and dynamic interaction between the real world and the computing system in real-time (Wolf, 2009). In this scenario, computing systems are commonly formed of many embedded units

that are heterogeneously networked. One of the challenges facing future systems is to address the separation of modal control design from the deployment of executable code. In such systems, it would be difficult, if not impossible, to analyze off-line all the possible combinations of processor loads. Because of the large range of control process domains it is difficult to address control modality only through the parameterization of pre-loaded local controllers. For this reason, it is worthwhile attempting to define new flexible architectures that enable computing systems to optimally adapt to potential changes in the

[☆]This work has been supported by the Spanish government as part of the SIDIRELI project (DPI2008-06737-C02-02).

*Tel: +34 665368512. Fax: +34 963877579

**Tel: +34 963877000 Ext:75706. Fax: +34 963877579

Email addresses: jacopa@ai2.upv.es (Javier O. Coronel),

jsimo@disca.upv.es (José Simó Ten)

environment in a dynamic manner.

From the point of view of embedded control systems, it is well known that the integration of controller design and real-time scheduling is necessary (Cervin, 2003). However, this approach is not enough in dynamic environments with limited computing resources. Therefore, additional mechanisms must be included in the design to optimize the use of resources and provide adaptation to the changing environmental conditions. Likewise, specific quality of services (QoS) levels must be guaranteed to ensure correct system operation. These QoS levels maintain a suitable control performance and a temporal predictability in the control system.

In this context, changes in the system operation mode, or in the environment, may force the disabling of some controllers and the enabling of others. In addition, controller switching mechanisms must be added in both local and distributed forms, resulting in workload changes for the embedded and networked units (Emberston & Bate, 2007; Real & Crespo, 2004). This can have a significant impact on the overall performance and real-time constraints of the system. For this reason, task and component migration approaches (Briao et al., 2007) may contribute to on-line redistribution and reallocation of workload – according to each situation.

An example of this kind of architecture is proposed in Simarro et al. (2008), which is developed in the context of a distributed real-time control system and from the perspective of control kernel middleware (Crespo et al., 2006; Albertos et al., 2006). The proposal includes features such as control basis, controller switching, and code delegation. From the point of view of CPS, these approaches partially bridge the abstraction gap (Lee, 2006).

In summary, cyber-physical systems should be able to deliver new levels of performance and efficiency thanks to sophisticated control computing co-design. Control systems mean that we must change our understanding of computing systems and accept that cyber-physical systems actively engage with the real world and real-time restrictions – and so expend real energy.

This paper explores methods that can be used in task migration when combined with techniques of energy management. These methods can maximize overall energy savings; facilitate thermal chip management by moving tasks away from the hot processing unit; balance the workload or concentration of parallel processing elements; decrease communication among those tasks that are particularly energy-efficient on wireless sensor networks (WSN) systems (Yi et al., 2009; Yuan & Wang, 2008; Kumar & Manimaran, 2007); and help guarantee the fulfillment of real-time system constraints.

Several power-aware techniques have been widely addressed in real-time literature (Piao et al., 2009; Tavares et al., 2008; Aydin et al., 2006; Pillai & Shin, 2001). These papers have discussed dynamic voltage scaling (DVS) of the processor, also known as dynamic voltage and frequency scaling. This approach exploits the convex, and normally quadratic relationship between CPU energy consumption and voltage. Additionally, these techniques have been extended to reduce the energy consumed during memory cycles (Cho & Chang, 2006; Liang et al., 2008) and network energy consumption (Yi et al., 2009; Kumar et al., 2008).

Let's consider a dynamic environment where an embedded and networked system operates with the support of task migration and processor frequency scaling. Assuming that the system knows where and when it must allocate tasks (Briao et al., 2007; Emberston & Bate, 2007), we must perform feasibility analyses when each task arrives and departs. This guarantees that the temporal requirements of the system will be accomplished during the task allocation phase. Additionally, a new processor speed (frequency scaling) should be also computed to enable the system to adapt itself to the new computational workload and so reduce energy consumption. Although some authors have carried out these two phases (feasibility analysis and frequency scaling computation) separately (AlEnawy & Aydin, 2005), these analyses are strongly related and in some cases can be performed together.

This work mainly focuses on a **proposal for a new algorithm** to be used on-line during the **task allocation and pro-**

cessor speed assignment phases. The algorithm uses fixed priority scheduling schemes with deadlines less than, or equal to, the period of the tasks when the dynamic processor workloads are being handled.

1.1. Related work

Few works cover task migration in the context of embedded and networked systems. Moreover, most algorithms are aimed at multiprocessor systems with soft real-time constraints (Yazdi et al., 2008; Emberson & Bate, 2007; Anderson et al., 2005). Some papers have addressed both task migration and energy consumption minimization (Briao et al., 2007; Chen, 2005). Briao et al. (2007) analyze the impact of task migration and justify its use because it compensates for the performance and energy costs involved in the system.

Several heuristics have been proposed for task allocation problems with deadlines equal to the periods (Mejia-Alvarez et al., 2004; Chen, 2007) when earliest-deadline-first (EDF) scheduling is in use. AlEnawy & Aydin (2005) use algorithms for allocating real-time tasks by applying rate monotonic (RM) scheduling. In their article, a comparison of some admission control algorithms is presented in function of complexity, feasibility, and energy consumption parameters. However, only tasks with deadlines that are equal to the period are analyzed. Moreover, aspects such as predictability of execution, behaviour related to the arrival of tasks, and the real computing cost of the algorithms are not taken into account.

At the CPU-level, DVS techniques can be classified as *static* or *dynamic*. Static algorithms for hard real-time systems use parameters such as period or minimum inter-arrival time, and assume that each task executes its worst-case execution time when selecting the processor frequency, and that this is statically decided before execution. Dynamic algorithms are based on the reclamation of additional slack resulting from the early completions of tasks. These are then used to further reduce the processor frequency and save more energy. These algorithms are applied at run-time.

By using static algorithms we can obtain a single proces-

sor frequency that never changes, or obtain variable frequencies that are statically decided before execution. An example of the former is Pillai & Shin (2001) in which the authors derive the minimum speed for a schedulable task set under EDF and propose an approximated method under RM. In Saewong & Rajkumar (2003) an algorithm that chooses a single frequency for a fixed priority scheduling scheme is proposed. In Bini et al. (2005) the authors present a method for approximating any speed level with two given discrete values that are switched as a pulse width modulation signal in order to obtain the average value. For the latter, in Mejia-Alvarez et al. (2004), and Saewong & Rajkumar (2003), the authors propose an approach whereby each task is assigned a different frequency. Several authors have published works that find the optimal voltage schedule for recurrent tasks – examples for periodic tasks include: Liu & Mok (2003), Yun & Kim (2003); and examples for periodic and aperiodic tasks include Zhong et al. (2007), Scordino & Lipari (2006). Furthermore, speed function and the characterization of the points in time at which speed changes occur has been presented in Liu & Mok (2003) and Gaujal & Navet (2007). However, a drawback in the works with statically established variable frequencies can appear if a task activation is lost or delayed. The drawback is that the entire frequency assignment will be affected, and this leads to missed deadlines.

Some authors, such as Piao et al. (2009), have proposed dynamic algorithms working in combination with static methods. These algorithms can be divided into inter-task and intra-task methods. In the inter-task algorithms, the processor frequency is determined task-by-task, whereas intra-task algorithms may adjust the frequency within the boundaries of a given task. In Kim et al. (2002), a performance comparison of several dynamic techniques is presented. In Pillai & Shin (2001), the authors propose dynamic algorithms for the EDF and RM schedulers. Saewong & Rajkumar (2003), Quan (2004), and Leung (2005) propose a simple dynamic scheme for fixed priorities. Aydin et al. (2004) presents a generic dynamic reclaiming algorithm, as well as an adaptive and speculative speed adjustment mechanism. Dynamic frequency changes for periodic and aperi-

riodic tasks are discussed in Piao et al. (2009).

The use of elastic scheduling to improve DVS management has been proposed in Marinoni & Buttazzo (2007) and in Zhu et al. (2004). Xia (2008) proposes energy management based on a feedback control scheduling methodology. The processor DVS analysis has been extended by other authors to reduce energy consumption in the memory (Cho & Chang, 2006; Liang et al., 2008) and network levels (Yi et al., 2009; Kumar et al., 2008).

However, some of the works mentioned above are based on an analysis of the hyper period. This approach can result in hard computing and is unsuitable for execution in run-time during task migration. Some other works consider only deadlines equal to periods and although this can be considered in specific cases, it can result in hard simplifications being applied in embedded control systems. Additionally, some studies have applied heuristics (Jejurikar & Gupta, 2004; Mejia-Alvarez et al., 2002) to obtain solutions that are sometimes far from optimal.

1.2. Contributions and paper organization

Global energy consumption in the system can contribute to the load balance criteria in the code deployment phase. This criterion can be combined with others such as schedulability, communication delays, control application correctness, and stability to determine the dynamic code movement and on-line load balancing in a system. However, we focus on the problem of task allocation, and more precisely on the feasibility analysis algorithms performed at task arrival or departure, as well as in energy management.

In this paper, we present novel algorithm that perform feasibility analyses and compute new processor frequencies when set tasks arrive and/or depart. Through extensive simulations, we evaluate the performance of this algorithm against other existing feasibility tests that have been adapted to compute the minimum processor frequency. This minimum frequency is computed in terms of energy consumption, acceptability ratio, and real computing costs. In addition, predictability in the execution and behaviour of the algorithms in relation to the contin-

uing arrival of tasks is analyzed. These terms will be explained later, particularly the real computing cost of algorithms. Our algorithms are based on DVS static algorithms and search for a single processor frequency while using a fixed priority scheduling scheme.

This paper is organized as follows. Section 2 describes the task model, processor model, and the evaluation of computational costs. In Section 3 we introduce energy minimization and the task acceptability problem in a dynamic environment with variable loads. In Section 4 an overview of a schedulability method-based feasibility region is presented. In Section 5 the adaptation of several existing feasibility tests for computing the processor scaling factor is shown. Additionally, in this section a **proposed** $\mathcal{A}$ approach is presented. Section 7 presents the results of the simulations used to evaluate the performance of several approaches when computing the α factor. And finally, Section 8 states our conclusions and future work.

2. System Model

2.1. Task model

In this paper we use a periodic task model. Let $\mathcal{T} = \{\tau_1, \tau_2, \dots, \tau_n\}$ denote a task set of n real-time preemptible tasks running on a uniprocessor system. Each task $\tau_i(T_i, D_i, C_i)$ is characterized by a period T_i , a relative deadline D_i , and a worst-case execution time C_i . We assume critical instants of activations for each task. Tasks are scheduled by a fixed priority scheduler (RM or DM) and ordered from highest to lowest priority, meaning that τ_1 has the highest priority and τ_n has the lowest priority.

We consider that only real-time tasks are executed in the CPU, this point being a simplification for the analysis rather than a restriction.

2.2. Processor model

The power consumption ($\mathcal{P}$) per CPU cycle is proportional to $C_L \cdot V_{dd}^2 \cdot f$, where C_L is the average load capacity, V_{dd} is the supply voltage of the processor, and f is the operating frequency of the processor. The maximum operating frequency

is a direct consequence of the supply voltage given by $f = K \frac{(V_{dd} - V_{th})^\sigma}{V_{dd}}$ where K is a constant specific for a given technology, V_{th} is the threshold voltage, and σ is the velocity saturation index where $1 \leq \sigma \leq 2$ (Saewong & Rajkumar, 2003; Pouwelse et al., 2001). However, the exact form of the power and frequency relation specifically depends on the processor hardware, and can be expressed in terms of CPU speed as a second or third degree polynomial function (Li & Ding, 2001; Aydin et al., 2004; Zhong et al., 2007; Marinoni & Buttazzo, 2007). In this article, the curve of function $\mathcal{P}$ (see Figure 1) was obtained for the relationship between power and speed on the Intel XScale processor (Zhong et al., 2007), while taking into account the processor consumption in idle state.

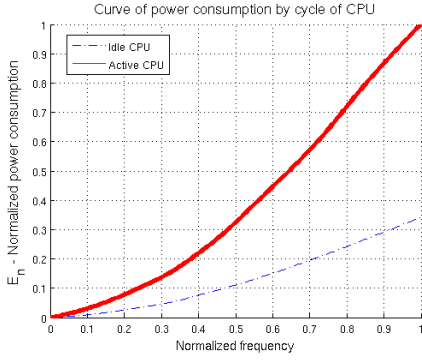


Figure 1: Normalized reference model of power consumption by CPU cycle as used in this paper.

For the analysis of the behaviour of several DVS (dynamic voltage scaling) algorithms we assume that the processor voltage can be changed continuously, and therefore, also the frequency. Considering the fact that most commercially available processors only provide a limited number of voltage levels, researchers have proposed algorithms to map continuous voltage levels to discrete levels (Bini et al., 2005). We therefore do not explore the effect of a limited number of processor frequencies. However, thanks to advances in power-supply electronics and CPU design, (Duan & Khatri, 2006), systems are increasingly able to operate using a wider voltage spectrum.

We also assume that the operating frequency will be adjusted by a normalized scaling factor α ($0 < \alpha \leq 1$), which will

be the equivalent to C scaled by a factor $1/\alpha$ and not affect the period nor the task deadlines. When α is equal to 0, the processor is in turn-off mode, and when α is 1, the processor runs at maximum clock frequency.

2.3. Evaluation of costs

To compare different proposals, some authors such as Bini & Buttazzo (2004) and Bini et al. (2003) have proposed the use of schedulability complexity tests based on the number of steps and iterations. However, if these comparisons are made from the viewpoint of computational cycles, the results could be strongly altered. For example, the cost of 100 loop sums is not the same as 100 loop divisions. Therefore, the complexity of algorithms to calculate a suitable α on embedded systems should be evaluated by the number of machine cycles required for each mathematical operation: especially on-line executions. In most embedded systems, it is well known that division operations have a higher computational cost than other operations. The routines of integer division using the Newton-Raphson approach can last between 20 and 100 cycles (Sloss et al., 2004), depending on the implementation and range of input operands. Therefore, these aspects will be considered in the evaluation of algorithms.

3. Scaling Frequency and Energy Consumption with Dynamic Loads

Energy consumption has been defined as an integral over the time t of $\mathcal{P}$:

$$E(S_y, t) = \int_0^{t_a} \mathcal{P}(F(t, \mathcal{T}(t))) dt$$

where S_y identifies the type of scheduling, and suffix y can be FP or DP for fixed and dynamic priority scheduling respectively. $\mathcal{P}$ is the power consumed by the CPU cycle and this depends on the CPU frequency function. The clock frequency (F) is composed as a function of time t and $\mathcal{T}$ task sets for each time instant t . Figure 2 shows an example of the scaling of processor frequency as a function of time and task set. In this

figure, two different task sets are represented over time. Obviously, the profile frequency scaling factor $\alpha(t)$ should change just when the processor changes the task set. This is because each α is calculated for a given task set.

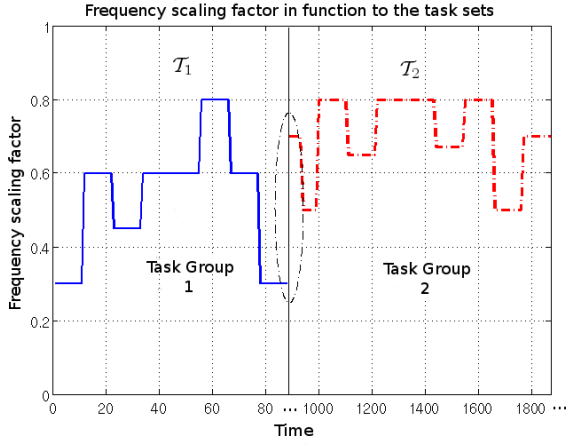


Figure 2: Example of the scaling of processor frequency α as a function of time and task set.

The switching between task sets is represented in Figure 2 as a dashed ellipse. The resultant task set is not known in advance because these are dynamically defined according to code delegate (J.L.Posadas et al., 2008; Emberson & Bate, 2007). Nevertheless, existing mode change protocols for real-time systems (Real & Crespo, 2004) can be adjusted and applied to the movement of real-time components. These change the task set demand by:

- repeating the schedulability analysis for the whole real-time system,
- recalculating the function $\alpha(t)$ for scaling the frequency of the processor.

These analyses can be performed together. One important aspect to consider in task migration is the computational cost involved in recalculating these parameters. This is because this calculation may mean an **increase in the energy consumption** and also an **increase in the time** to perform the incorporation or rejection of tasks.

After a migration of components, a *resultant task set* ($\mathcal{T}_{x+1}$) will be understood as a modification of the current task set ($\mathcal{T}_x$)

due to inclusions and expulsions of tasks. Therefore, the $\mathcal{T}_{x+1}$ set will be defined as:

$$\mathcal{T}_{x+1}^N = \mathcal{T}_x^n + \mathcal{T}_I^m - \mathcal{T}_E^p = \mathcal{T}_x^n + \Delta\mathcal{T}$$

where the size of $\mathcal{T}_{x+1}$ is $N = n + m - p$. $\mathcal{T}_I$ is the task set that migrates to the processor, and $\mathcal{T}_E$ is the task set that migrates from the processor, where $\mathcal{T}_E \subseteq \mathcal{T}_x$.

The energy consumption function when assuming the mobility of components at instants of time t_x will be given by:

$$E(S_y, t_a) = \int_{t_0}^{t_1} \mathcal{P}(F(t, \mathcal{T}_0))dt + \dots + \int_{t_x}^{t_{x+1}} \mathcal{P}(F(t, \mathcal{T}_x))dt + \dots + \int_{t_{x+k}}^{t_a} \mathcal{P}(F(t, \mathcal{T}_{x+k}))dt \quad (1)$$

where $F(t, \mathcal{T}_x)$ is the function of processor speed obtained from the task set $\mathcal{T}_x$ during a time period.

4. An Overview of a Schedulability-Analysis Based Feasibility Region

In this article the schedulability analysis and static minimum solutions to α are addressed from the perspective of an analysis in C-space. In this space, the task computation times C_i are considered as variable parameters, whereas the periods and the deadlines are fixed parameters. Consequently, we obtain a constraint on the C_i variables, which is a function of all T_i and D_i .

The analysis is reduced to verify if a point is inside a region delimited by coordinates C_i . This region is known as $\mathbf{R}_n$, and was originally proposed in Lehoczy et al. (1989).

The formal formulation of the feasibility region $\mathbf{R}_n$ was derived from Lehoczy et al. (1989) and conveniently demonstrated in Bini & Buttazzo (2004) with the following theorem:

Theorem 1. (Theorem 2 in Bini & Buttazzo (2004)). *When deadlines D_i are equal to periods T_i , the region $\mathbf{R}_n$ of a schedulable task set is given by:*

$$\mathbf{R}_n(T_1, \dots, T_n, D_1, \dots, D_n)|_{D_i=T_i} = \{(C_1, \dots, C_n) \in \mathbb{R}_+^n :$$

$$\max_{i=1 \dots n} \min_{t \in S_i} \sum_{j=1}^i \left\lceil \frac{t}{T_j} \right\rceil C_j \leq t\} \quad (2)$$

where $\mathcal{S}_i = \{kT_j : j = 1 \dots i, k = 1 \dots \lfloor \frac{T_i}{T_j} \rfloor\}$.

The elements of $\mathcal{S}_i$ conceptually represent the arrival times of tasks with a priority higher than τ_i before deadlines D_i and D of task τ_i , and assuming $\phi_i = 0$. This subset of values is called *rate monotonic scheduling points* $\mathcal{S}$. In Bini & Buttazzo (2004), this theorem is formulated through logical operators to represent the max ($\wedge$) and min ($\vee$).

For a better understanding and characterization in terms of complexity and implementation costs of this approach, we will apply the following notation for a matrix representation of Equation 2. The element set $\mathcal{S}_i$ is represented by the matrix:

$$\mathcal{S}_i = \{s_{kl}\}_{k \in [1, \dots, i] \wedge l \in [1, \dots, \lfloor \frac{T_i}{T_1} \rfloor]} \quad (3)$$

where s_{kl} is equal to the product:

$$s_{kl} = \begin{bmatrix} T_1 \\ T_2 \\ \dots \\ T_k \end{bmatrix} \otimes \begin{bmatrix} 1 & 2 & 3 & \dots & \lfloor \frac{T_i}{T_k} \rfloor \end{bmatrix} = \begin{bmatrix} s_{11} & s_{12} & \dots & s_{1l} \\ s_{21} & s_{22} & \dots & s_{2l} \\ \dots & \dots & \dots & \dots \\ s_{k1} & 0 & \dots & 0 \end{bmatrix}$$

The column numbers in each row of the matrix s_{kl} are determined by the difference between period tasks ($\lfloor \frac{T_i}{T_k} \rfloor$), thus several elements from the matrix could be 0, and as a result, the set $\mathcal{S}_i$ can be defined as $\mathcal{S}_i^* = \mathcal{S}_i / s_{kl} \neq 0$.

The total number of elements of $\mathcal{S}_i^*$ is determined by:

$$\text{Size}(\mathcal{S}_i^*) = \sum_{j=1}^i \left\lfloor \frac{T_i}{T_j} \right\rfloor \quad (4)$$

Moreover, Equation 2 can be expressed again as:

$$\max_{i=1 \dots n} \min M_i(\mathcal{S}_i) \leq \mathcal{S}_i \quad (5)$$

where $M_i(\mathcal{S}_i)$ is:

$$M_i(\mathcal{S}_i) = \{m_{kl}\}_i \quad (6)$$

$$\{m_{kl}\}_i = \sum_{j=1}^i \left\lfloor \frac{s_{kl}}{T_j} \right\rfloor C_j \quad / \quad s_{kl} \in \mathcal{S}_i^*$$

where m_{kl} is a matrix with the same dimension as the matrix s_{kl} . This will result in a matrix sum:

$$\{m_{kl}\}_i = \begin{bmatrix} \left\lfloor \frac{s_{11}}{T_1} \right\rfloor C_1 & \dots & \left\lfloor \frac{s_{1l}}{T_1} \right\rfloor C_1 \\ \left\lfloor \frac{s_{21}}{T_1} \right\rfloor C_1 & \dots & \left\lfloor \frac{s_{2l}}{T_1} \right\rfloor C_1 \\ \dots & \dots & \dots \\ \left\lfloor \frac{s_{k1}}{T_1} \right\rfloor C_1 & \dots & \left\lfloor \frac{s_{kl}}{T_1} \right\rfloor C_1 \end{bmatrix} + \dots +$$

$$\begin{bmatrix} \left\lfloor \frac{s_{11}}{T_1} \right\rfloor C_i & \dots & \left\lfloor \frac{s_{1l}}{T_1} \right\rfloor C_i \\ \left\lfloor \frac{s_{21}}{T_i} \right\rfloor C_i & \dots & \left\lfloor \frac{s_{2l}}{T_i} \right\rfloor C_i \\ \dots & \dots & \dots \\ \left\lfloor \frac{s_{k1}}{T_i} \right\rfloor C_i & \dots & \left\lfloor \frac{s_{kl}}{T_i} \right\rfloor C_i \end{bmatrix}$$

From this expression, it is possible to determine for which repeated values s_{kl} gives redundant values of m_{kl} as the result. Therefore, $\mathcal{S}_i^*$ can be redefined as:

$$\mathcal{S}_i^* = \mathcal{S}_i / s_{kl} \neq 0 \wedge \exists! s_{kl}$$

Region $\mathbf{R}_n$ (equation 2) can be expressed as:

$$\mathbf{R}_n(T_1, \dots, T_n, D_1, \dots, D_n)_{|D_i=T_i} =$$

$$\{(C_1, \dots, C_n) \in \mathbb{R}_+^n : \max_{i=1 \dots n} \min M_i(\mathcal{S}_i^*) \leq \mathcal{S}_i^*\} \quad (7)$$

where $M_i(\mathcal{S}_i^*)$ is defined in Equation 6. Region $\mathbf{R}_n$ is delimited by planes which define the space where the points C_i must be located. However, for analysis of large task sets, the number of equations to be checked is huge and equals the sum of the number of elements in all $\mathcal{S}_i$ (such as shown in Equation 4). This prevents a practical on-line application of Theorem 1. Nevertheless, such as is demonstrated in Bini & Buttazzo (2004), solving Equation 7 results in many equations that are useless, so the idea is to reduce the number of equations by eliminating the redundant elements in $\mathcal{S}_i$. This approach has led to the formulation of the following theorem:

Theorem 2. (Theorem 3 in Bini & Buttazzo (2004)). The region of the schedulable task sets $\mathbf{R}_n$, such as defined by 2, is given by:

$$\mathbf{R}_n(T_1, \dots, T_n, D_1, \dots, D_n) = \{(C_1, \dots, C_n) \in \mathbb{R}_+^n :$$

$$\max_{i=1 \dots n} \min_{t \in \mathcal{P}_{i-1}(D_i)} C_i + \sum_{j=1}^{i-1} \left\lfloor \frac{t}{T_j} \right\rfloor C_j \leq t\} \quad (8)$$

where $\mathcal{P}_i(t)$ is defined by the following expression:

$$\begin{cases} \mathcal{P}_0(t) = \{t\} \\ \mathcal{P}_i(t) = \mathcal{P}_{i-1}(\lfloor \frac{t}{T_i} \rfloor T_i) \cup \mathcal{P}_{i-1}(t) \end{cases} \quad (9)$$

The total number of elements of $\mathcal{P}_{i-1}(t)$ is determined by:

$$\text{Size}(\mathcal{P}_{i-1}) = \sum_{n=1}^i 2^{n-1} \quad (10)$$

Note that the difference between this theorem and Theorem 2 is only the presence of the $\mathcal{P}_{i-1}(t)$ set, instead of $\mathcal{S}_i$. $\mathcal{P}_{i-1}(t)$ is a $\mathcal{P}_{i-1}(t) \subseteq \mathcal{S}_i$, and this strongly reduces the number of equations to be evaluated. As a consequence, the time needed to establish whether a set task is inside the region $\mathbf{R}_n$ is limited.

Following the same expression as Equation 5, a periodic task set is feasibly schedulable using Theorem 2 if and only if:

$$\max_{i=1 \dots n} \min M_i(\mathcal{P}_{i-1}) \leq \mathcal{P}_{i-1}^* \quad (11)$$

where $\mathcal{P}_{i-1}^* = \mathcal{P}_{i-1} / p \exists! p$.

5. Computing the Processor Scaling Factor

In this section, we state – as based on the test described above and other schedulability tests – that it is possible to compute the frequency scaling factor α for task sets with $D=T$ and $D \leq T$. Furthermore, we establish that if this same computing approach can be used for analyzing the schedulability of periodic task sets.

Not all execution cycles are scaled for the processor speed because some operations deal with memory, or other I/O devices, whose access time may be fixed (Bini et al., 2005). However, access to memory can also be performed at different speeds (Marvell, 2008), and so we should use two scaling factors: α to scale the processor; and β to scale the system bus. To simplify the analysis, we assume β as a fixed parameter equal to 1, such that the worst-case execution time C can be expressed as a parameter split in two: a parameter C^f which is dependent on the processor frequency; and a fixed parameter C^m which is not.

Therefore, C_i can be expressed as the execution time at the maximum frequency of the processor with respect to the frequency scaling factor α plus a constant execution time:

$$C_i = \frac{C_i^f}{\alpha} + C_i^m \quad (12)$$

5.1. LL approach

Theorem 3. The computing of the frequency scaling factor α using the well-known approximate schedulability test of Liu Laylan (Sha et al., 2004) (LL) is given by:

$$\alpha_x^{LL} = \frac{U_f}{n(2^{1/n} - 1) - U_m} \quad (13)$$

where U_f is the sum of the whole utilization of the part of the task set that depends on the processor frequency, U_m is the utilization of the part of the task set that is independent of the processor.

After a migration of components, the new scaling factor LL α_{x+1}^{LL} is defined as:

$$\alpha_{x+1}^{LL} = \frac{(U_f^{\alpha_x}(t_x) + \Delta U_f^{\alpha_x}(t_{x+1}))\alpha_x}{n(2^{1/n} - 1) - U_m(t_x) - \Delta U_m(t_{x+1})} \quad (14)$$

where ΔU is the difference of utilization between task inclusion and expulsion. And α_x is the frequency scaling factor before the migration of components. However, it is well-known that the LL test is a sufficient but unnecessary test, and so the scaling factor α^{LL} may not be the minimum α that minimizes energy consumption. This test could be used only for a task set with $D=T$.

5.2. HB approach

Theorem 4. The frequency scaling factor based on the approximate hyperbolic bound (Bini et al. (2003)) test can be computed as:

$$\alpha_x^{HB} = \max\{\alpha_n\} \quad (15)$$

where α_n represent the n possible values that can be assigned to alpha, and which are the result of solving the roots of the product of all the n utilizations plus one $\left(\prod_{i=1}^n \left(\frac{U_i^f}{\alpha_x} + U_i^m + 1\right) - 2 = 0\right)$ of the task set on the computer system.

To obtain the scaling factor after a migration of components, it is necessary to solve the product by deducing α_{x+1} from the equation:

$$2 = \prod_{i=1}^n \left(\frac{U_i^f \cdot \alpha_x}{\alpha_{x+1}} + U_i^m + 1 \right) \cdot \frac{\prod_{j=1}^{inc} \left(\frac{U_j^f \cdot \alpha_x}{\alpha_{x+1}} + U_j^m + 1 \right)}{\prod_{k=1}^{exp} \left(\frac{U_k^f \cdot \alpha_x}{\alpha_{x+1}} + U_k^m + 1 \right)}$$

where *inc* and *exp* are respectively the number of tasks included and expelled to and from the current ($\mathcal{T}(t_x)$) task set. α_x is the frequency scaling factor before the migration of components. In the same way as the LL approach, the HB approach cannot be used for task sets where $D < T$.

5.3. LLM approach

As an extension of LL test, an utilization bound test for tasks with pre-period deadlines has been used Racu et al. (2005), and some changes were made to find an approximate value for α . We termed this the LLM test.

Theorem 5. *The frequency scaling factor based on an approximate utilization bounds test is given by:*

$$\alpha^{LLM} = \max_{i=1 \dots n} \frac{f_i^f}{U(p, \Delta_i) - f_i^m} \quad (16)$$

where f_i^f and f_i^m are respectively:

$$f_i^{f/m} = \sum_{j \in H_p} \frac{C_j^{f/m}}{T_j} + \sum_{k \in H_1} \frac{C_k^{f/m}}{T_i} + \frac{C_i^{f/m}}{T_i}$$

where H_1 consists of a set of higher priority tasks that preempt task τ_i only once before its deadline at D_i , and H_p consists of a set of higher priority tasks that may often preempt task τ_i before the deadline at D_i .

$U(p, \Delta_i)$ is:

$$U(p, \Delta_i) = \begin{cases} p((2\Delta_i)^{1/n} - 1) + 1 - \Delta_i & 0.5 \leq \Delta_i \leq 1 \\ \Delta_i & 0 \leq \Delta_i < 0.5 \end{cases}$$

where Δ_i and p are:

$$\Delta_i = \frac{D_i}{T_i} \quad p = \text{num}(H_p) + 1$$

5.4. RTA approach

As is well known, the response time analysis (RTA) (Sha et al., 2004) recurrent approach is an exact test, which we can change to find an exact value for α that minimizes energy consumption. It is based on Saewong & Rajkumar (2003) and we can obtain a static and exact solution to compute the frequency scaling factor for task sets with $D=T$ and $D < T$.

5.5. Approach based on the schedulability region ($\mathcal{P}$ and $\mathcal{S}$ approaches)

Taking into account the above regarding the representation of the feasibility condition as a region in the C-space, it is easy to think that C_i can be tuned so that its values are in the feasibility region. In Bini et al. (2006) this topic is addressed from the perspective of sensitivity analysis.

For our purpose, in Equation 11 and using Equation 12, C_i becomes the design variable and a constant parameter. The scaling factor α is the interesting component that results in the following theorem:

Theorem 6. *The static scaling factor of the clock frequency that minimizes CPU energy consumption while preventing missed deadlines when given a task set $\mathcal{T}$ is defined by:*

$$\alpha_{\min}^{\mathcal{P}}(\mathcal{T}) = \max_{i=1 \dots n} \min M_i^*(\mathcal{P}_{i-1}) \quad (17)$$

where $\mathcal{P}_{i-1}$ are the scheduling points defined by Equation 9.

This expression can also be used with the scheduling points $\mathcal{S}_i$:

$$\alpha_{\min}^{\mathcal{S}}(\mathcal{T}) = \max_{i=1 \dots n} \min M_i^*(\mathcal{S}_i) \quad (18)$$

where $M_i^*(\mathcal{S}_i)$ is given by:

$$M_i^*(\mathcal{S}_i) = \sum_{j=1}^i \frac{\left\lceil \frac{s_{kl}}{T_j} \right\rceil C_j^f}{t - \left\lfloor \frac{s_{kl}}{T_j} \right\rfloor C_j^m} \quad / \quad s_{kl} \in \mathcal{S}_i^*$$

s_{kl} is defined in Equation 3. C_j^f and C_j^m are defined in Equation 12.

Let the scaling factor be normalized between $0 < \alpha \leq 1$, the computer system is feasibly schedulable **if and only if** α is less than or equal to 1, otherwise, the processor is unable to schedule the task set even though it is running at full speed.

PROOF. Based on definitions in Equations 5 and 12, and following several changes to simplify the analysis, we have:

$$\begin{aligned} \max_{i=1 \dots n} \min M_i(\mathcal{S}_i) &\leq \mathcal{S}_i \\ \max_{i=1 \dots n} \min \sum_{j=1}^i \left\lceil \frac{s_{kl}}{T_j} \right\rceil C_j &\leq s_{kl} \quad / \quad s_{kl} \in \mathcal{S}_i^* \end{aligned}$$

$$\max_{i=1\dots n} \min \sum_{j=1}^i \left\lceil \frac{s_{kl}}{T_j} \right\rceil \left(\frac{C_j^f}{\alpha} + C_j^m \right) \leq s_{kl} \quad / \quad s_{kl} \in \mathcal{S}_i^*$$

$$\max_{i=1\dots n} \min \sum_{j=1}^i \left\lceil \frac{s_{kl}}{T_j} \right\rceil \frac{C_j^f}{\alpha} \leq s_{kl} - \max_{i=1\dots n} \min \sum_{j=1}^i \left\lceil \frac{s_{kl}}{T_j} \right\rceil C_j^m$$

deducing α from the expression, we have:

$$\max_{i=1\dots n} \min M_i^*(\mathcal{S}_i) \leq \alpha \triangleq \alpha_{min}$$

where the modified matrix M (M_i^*) will be:

$$M_i^*(\mathcal{S}_i) = \{m_{kl}^*\}_i \quad (19)$$

$$/ \quad \{m_{kl}^*\}_i = \sum_{j=1}^i \frac{\left\lceil \frac{s_{kl}}{T_j} \right\rceil C_j^f}{s_{kl} - \left\lceil \frac{s_{kl}}{T_j} \right\rceil C_j^m} \quad / \quad s_{kl} \in \mathcal{S}_i^*$$

These expressions can be expanded to the points $\mathcal{P}_{i-1}$.

6. Proposed approach $\mathcal{A}$

To bound the schedulability region for periodic fixed priority systems the analysis in the C-space is mainly based on the point set $\mathcal{S}_i$ or $\mathcal{P}_i$ (theorem 1 and 2). The number of mathematical operations and the accuracy of this analysis is determined by the structure and the number of scheduling points.

In Section 4, two expressions for calculating the schedulability region were described (Equations 7 and 11). These regions are delimited by two different point sets ($\mathcal{S}_i$ and $\mathcal{P}_{i-1}$), which differ in size: $Size(\mathcal{S}_i) \gg Size(\mathcal{P}_{i-1})$ (equations 4 and 10). The size of point set $\mathcal{S}$ can be much greater than the size of $\mathcal{P}$. However, both approaches are equally precise, and both are sufficient and necessary tests.

6.1. Reduced approach using the schedulability region

In this section we propose a reduction in the number of scheduling points required to compute the schedulability region. The objective of this simplification is to substantially reduce the computational cost of the DVS algorithm, but at the same time, preserve the accuracy of the schedulability analysis.

The reduction in these scheduling points will be based on points $\mathcal{P}_{i-1}$, which in turn are a reduction in the point set $\mathcal{S}_i$

($\mathcal{P}_{i-1} \subseteq \mathcal{S}_i$). The new point set we call $\mathcal{A}_i$. This $\mathcal{A}_i$ set is a $\mathcal{A}_i \subseteq \mathcal{P}_{i-1}$ and a $\mathcal{A}_i \subseteq \mathcal{S}_i$.

The frequency scaling factor that uses the new set of scheduling points is defined by the following theorem:

Theorem 7. *The minimum scaling factor of the processor frequency that minimizes the energy consumption and guarantees no missed deadline given a task set $\mathcal{T}$, is defined as:*

$$\alpha_{min}^{\mathcal{A}}(\mathcal{T}) = \max_{i=1\dots n} \min M_i^*(\mathcal{A}_i) \quad (20)$$

where M_i^* is defined in Equation 19, but applied to the points set $\mathcal{A}$. This points set $\mathcal{A}_i$ is equal to:

$$\mathcal{A}_i(D_i) = \{D_i \cup \text{sched}\mathcal{A}_i(D_i)\} \quad (21)$$

where $\text{sched}\mathcal{A}(D_i)$ is the characteristic matrix in the computation of the points $\mathcal{A}_i$. $\text{sched}\mathcal{A}(D_i)$ is the following matrix expression:

$$\begin{bmatrix} \left\lceil \frac{D_i}{T_1} \right\rceil T_1 & \left\lceil \frac{D_i}{T_2} \right\rceil T_2 & \dots & \left\lceil \frac{D_i}{T_j} \right\rceil T_j \\ & \left\lceil \frac{D_i}{T_2} \right\rceil \frac{T_2}{T_1} T_1 & \dots & \left\lceil \frac{D_i}{T_j} \right\rceil \frac{T_j}{T_{j-1}} T_{j-1} \\ & & \dots & \left\lceil \frac{D_i}{T_j} \right\rceil \frac{T_j}{T_{j-1}} \frac{T_{j-1}}{T_{j-2}} T_{j-2} \\ & & & \dots \dots \\ & & & \left\lceil \frac{D_i}{T_j} \right\rceil \frac{T_j}{T_{j-1}} \frac{T_{j-1}}{T_{j-2}} \dots \frac{T_{j-k-1}}{T_{j-k}} T_{j-k} \end{bmatrix}$$

$$\forall j \in [1, \dots, i-1] \wedge \forall k \in [0, \dots, i-2] \quad (22)$$

Where the dimension of this matrix is given by $(i-1) \times (i-1)$.

Taking into account that the scaling factor has been normalized between $0 < \alpha \leq 1$, the whole system will be schedulable if α is less than or equal to 1. If α is greater than 1, there is a low probability that the system can be really schedulable, on the contrary, if α is less than or equal to 1, the system is always schedulable.

PROOF. Schedulability region analysis (Bini & Buttazzo, 2004; Lehoczky et al., 1989) is mainly based on a worst-case workload analysis of the task set. Using the concept of workload, the schedulability condition of a particular task τ_i can be expressed by:

$$C_i + W_{i-1}(D_i) \leq D_i \quad (23)$$

The processor demand in the interval $[0, t]$ is defined by $\sum_{j=1}^i \left\lceil \frac{t}{T_j} \right\rceil C_j$. To be a feasible schedule, the workload in $[t, D]$ should be smaller than the length of the interval, therefore:

$$\forall t \in [0, D_i] \quad W_i(D_i) - W_i(t) \leq (D_i - t)$$

The latter equation can also be expressed as follows:

$$\forall t \in [0, D_i] \quad W_i(D_i) \leq \sum_{j=1}^i \left\lceil \frac{t}{T_j} \right\rceil C_j + (D_i - t) \quad (24)$$

Moreover, we define a term $\varphi_i(D_i)$ as the last instant in $[0, D_i]$ during which the processor is idle:

$$\varphi_i(D_i) = \max\{t \in [0, D_i] \mid t \notin \text{active tasks in instants } t\}$$

From this expression we can induce that the processor is always busy in $[\varphi_i, D_i]$, and thus the workload of the i highest priority tasks during such interval is $(D_i - \varphi_i(D_i))$. The definition of φ is needed because it can be useful for simplifying the computation of $W_i(D_i)$. Hence,

$$W_i(D_i) = \sum_{j=1}^i \left\lceil \frac{\varphi_i(D_i)}{T_j} \right\rceil C_j + (D_i - \varphi_i(D_i)) \quad (25)$$

However, using this expression we move from the complexity of the workload estimation itself (using continuous interval $[0, D_i]$) to the complexity of the $\varphi_i(D_i)$ search. Nevertheless, a set of possible values of $\varphi_i(D_i)$ can be restricted. With this aim, an initial set of possible values can be comprised for the deadline of task i and the arrival times of tasks with a priority higher than τ_i before D_i . This is because the last instant of time in which the processor can be idle (φ_i) matches exactly with the arrival times of the tasks. This set is defined by the point set S_i described in Section 4, such that $\varphi_i(t) \in S_i(t)$. This set of possible values of S_i can be further reduced by the set of values proposed in Equation 9, such that $\varphi_i(t) \in \mathcal{P}_i(t)$. From the evaluation of the schedulability points in Expression 24, the minimum value corresponds to the workload i and is the same when $t = \varphi_i(D)$ (Bini & Buttazzo, 2004), that is:

$$W_i(D_i) = \min_{t \in S_i(D_i) \cup \mathcal{P}_{i-1}(D_i)} \sum_{j=1}^i \left\lceil \frac{t}{T_j} \right\rceil C_j + (D_i - t) \quad (26)$$

Using the equation in 23 for a whole task set, we obtain Equations 2 and 8.

Having established the principle of the schedulability region and feasibility conditions, we focus on proof of the reduced set $\mathcal{A}$. To obtain the reduced point set $\mathcal{A}$ based on the points $\mathcal{P}$, we propose to extract a characteristic function of the points $\mathcal{P}_{i-1}$, in order to obtain the most important points that enable narrowing and delineating the area where the points $\mathcal{P}_{i-1}$ are located at the time. This characteristic function is defined by means of a matrix expression, which is given in Equation 22. The set $\mathcal{A}_i$ is given by:

$$\mathcal{A}_i(D_i) = \{D_i \cup \text{sched} \mathcal{A}(D_i)\}$$

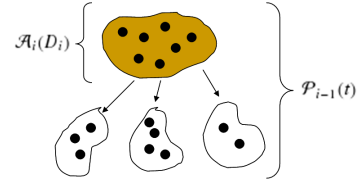


Figure 3: The set $\mathcal{P}_{i-1}(D_i)$ can be obtained from the set $\mathcal{A}_i$

From the set $\mathcal{A}(D_i)$, we can obtain additional point subsets that in total constitute the set $\mathcal{P}_{i-1}$ (see Figure 3). The number of elements of $\mathcal{P}_{i-1}(D_i)$ and $\mathcal{A}(D_i)$ are determined respectively by the equations 10 and 28. In Figure 4, the points of the sets $\mathcal{P}_{i-1}$ and $\mathcal{A}_i$ are compared for a specific set of 15 periodic tasks. Notice that the regions of point concentration $\mathcal{P}_{i-1}$ are delimited by the points obtained from the characteristic matrix $\mathcal{A}_i$.

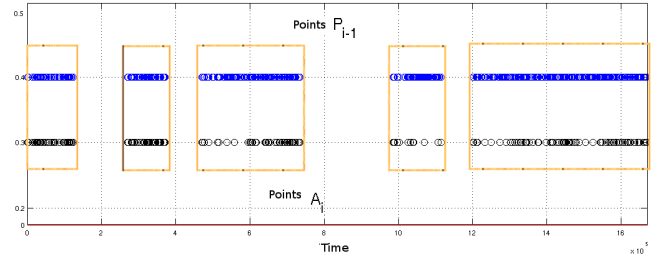


Figure 4: Comparison of the points $\mathcal{P}_{i-1}$ with the points $\mathcal{A}_i$

By definition the set $\mathcal{A}$ uses fewer points than $\mathcal{P}$ for the estimation of the values of $\varphi_i(D_i)$, therefore two things can happen:

1. The exact value of $\varphi_i(D_i)$ is within the values obtained in the calculation of $\mathcal{A}_i(D_i)$.
2. There is a difference between the exact value $\varphi_i(D_i)$ and the values obtained by means of $\mathcal{A}_i(D_i)$.

In the first case, we get an exact result for $W_i(D_i)$ and so the schedulability conditions exposed will be accurately tested. Moreover, we could also compute the frequency scaling factor α within a margin of error 0.

In the second case, let the smallest workload W_i correspond with the evaluation at $t = \varphi_i(D_i)$, and assuming that $\varphi_i(D_i)$ is not included in the points $\mathcal{A}_i(D_i)$, then W_i when $t \in \{\mathcal{A}_i(D_i)\}$ is always fulfilled that:

$$\varphi_i(D_i) \notin \mathcal{A}_i(D_i) \Rightarrow W_i(D_i)^{t=\mathcal{A}_i(D_i)} \geq W_i(D_i)^{t=\varphi_i(D_i)} \quad (27)$$

This proves that the result of calculating the frequency scaling factor using set $\mathcal{A}$ is always valid. Using the set $\mathcal{A}$ we will obtain an $\alpha^{\mathcal{A}}$ equal or greater than the exact value. However, we will never obtain a value below the exact scaling factor. This guarantees a result that is *sufficient*, although in some cases not *necessary*.

Based on the analysis of the schedulability region, we have:

$$\begin{aligned} \max_{i=1\dots n} \min M_i(\mathcal{A}_i) &\leq \mathcal{A}_i \\ \max_{i=1\dots n} \min \sum_{j=1}^i \left\lceil \frac{a}{T_j} \right\rceil C_j &\leq a \quad / \quad a \in \mathcal{A}_i^* \\ \max_{i=1\dots n} \min \sum_{j=1}^i \left\lceil \frac{a}{T_j} \right\rceil \left(\frac{C_j^f}{\alpha} + C_j^m \right) &\leq a \quad / \quad a \in \mathcal{A}_i^* \\ \max_{i=1\dots n} \min \sum_{j=1}^i \left\lceil \frac{a}{T_j} \right\rceil \frac{C_j^f}{\alpha} &\leq a - \max_{i=1\dots n} \min \sum_{j=1}^i \left\lceil \frac{a}{T_j} \right\rceil C_j^m \end{aligned}$$

where a is the element set that forms the set $\mathcal{A}_i^*$ (see Equation 21) and $\mathcal{A}_i^*$:

$$\mathcal{A}_i^* = \{a\} / a \in \mathcal{A}_i \wedge \exists! a \in \mathcal{A}_i$$

The total number of elements of $\mathcal{A}_i$ is determined by:

$$\text{Size}(\mathcal{A}_i) = i + \sum_{n=1}^i \frac{n(n-1)}{2} \quad (28)$$

Deducing the factor α from $\{m\}_i$, we have:

$$\max_{i=1\dots n} \min M_i^*(\mathcal{A}_i) \leq \alpha \triangleq \alpha_{\min}$$

where,

$$\begin{aligned} M_i^*(\mathcal{A}_i) &= \{m^*\}_i \quad / \\ / \{m^*\}_i &= \sum_{j=1}^i \frac{\left\lceil \frac{a}{T_j} \right\rceil C_j^f}{a - \left\lceil \frac{a}{T_j} \right\rceil C_j^m} \quad / \quad a \in \mathcal{A}_i^* \end{aligned}$$

We can obtain a maximum error in the computation of α using the set $\mathcal{A}_i$ when $\varphi_i(D_i) \notin \mathcal{A}_i(D_i)$, i.e. when the $\mathcal{A}_i(D_i)$ approach finds an approximate value (φ_i^*) equal to:

$$\text{error}_{\alpha^{\mathcal{A}}} = \sum_{j=1}^i \left(\left\lceil \frac{\varphi_i^*(D_i)}{T_j} \right\rceil \frac{1}{\varphi_i^*(D_i)} - \left\lceil \frac{\varphi_i(D_i)}{T_j} \right\rceil \frac{1}{\varphi_i(D_i)} \right) C_j \quad (29)$$

7. Experimental Evaluation when $D = T$ and $D \leq T$

We have performed extensive simulations of the algorithm proposed ($\mathcal{A}$ approach) and the other algorithms described ($\mathcal{P}$, RTA , LLM , HB and LL approaches) to experimentally evaluate the performance of the dynamic code delegation for task sets with $D = T$ and $D \leq T$.

We propose three group of random tasks: A, B and C. Group A consists of tasks with short periods, group B consists of tasks with middle periods, and group C consists of tasks with large periods. Each task group consists of 20 tasks. The range of periods for the task groups are shown in Table 1.

Group	Lower Bound	Upper Bound
A	2000 μs	40000 μs
B	40001 μs	600000 μs
C	600001 μs	4000000 μs

Table 1: Range of periods for task groups A, B and C.

Three types of arrival of tasks at the processor ($L11$, $L12$ and $L13$) have also been simulated. For $L11$ the highest priority tasks arrive first, for $L12$ the medium priority tasks arrive first, and for $L13$ the least priority tasks arrive first. The simulation consists of sets of 20 tasks.

Task groups A, B, and C are combined with arrival types $L11$, $L12$ and $L13$, and this results in nine different sets of tasks.

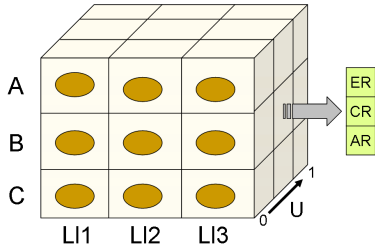


Figure 5: Test sets for analysing approaches for computing the scaling factor.

These, in turn, will be evaluated on five discrete utilizations U (0.3 - 0.5 - 0.7 - 0.8 - 0.95).

Each test is characterized by an acceptance ratio **AR**, computational cost **CR**, and energy consumption **ER** (see Figure 5). The results of these simulations will be presented on three components of reference:

- The *cost* of each algorithm will be counted in accordance with the number and type of operations. Each operation is characterized with a predetermined number of machine cycles based on the ARM architecture (Sloss et al., 2004).
- The component of *energy over-consumption* is one less than the energy consumption with respect to the consumption achieved with an exact algorithm.
- The *rejection ratio* refers to one less than the number of accepted task sets $\mathcal{T}_{pt}$ with respect to those accepted by a necessary and sufficient algorithm $\mathcal{T}_p$: $1 - (\mathcal{T}_{pt}/\mathcal{T}_p)$.

These components were chosen so that a greater magnitude on the axes can be understood as a worse behaviour for an algorithm. In the tests, the worst result will be plotted for each task group and type of arrival. The calculation of α using the $\mathcal{S}$ approach was not used due to the high computational cost involved.

7.1. Rejection ratio versus computational cost

Figures 6 and 7 show the simulation results of the percentage of rejected tasks in comparison with the computational cost of calculating algorithms.

In the figures, we can see that all approaches for the calculation of α , except for LLM, LL, and HB, have a rejection ratio

equal to 0%. The LL approach has a dispersed rejection ratio that increases with utilization and which stretches from 0% to about 35%. The HB approach has a slightly lower rejection percentage than LL. When $D \leq T$, the rejection ratio for the LLM approach is increased to about 70%.

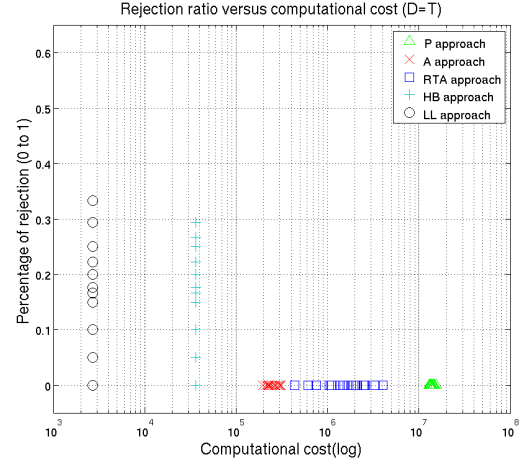


Figure 6: Rejection ratio versus computational cost when $D=T$. Note that the percentage of over-consumption has been plotted between 0 and 0.6.

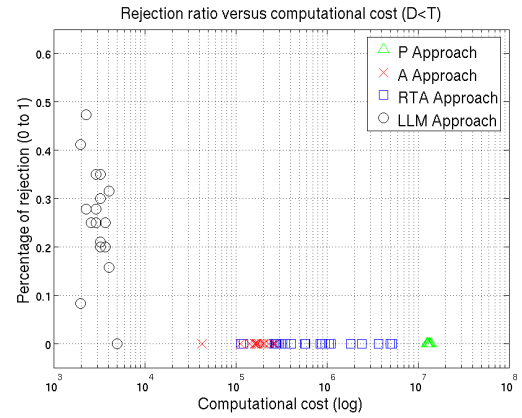


Figure 7: Rejection ratio versus computational cost when $D<T$.

From the perspective of cost, although the methods of LLM, LL, and HB show the largest rejection percentages, they also have lower costs. In both cases, $D = T$ and $D \leq T$, the computational cost of using the RTA approach is dispersed and unpredictable, and their costs are sometimes among the highest of all the methods, exceeded only by the $\mathcal{P}$ approach. The cost is sometimes less than the proposed $\mathcal{A}$ approach.

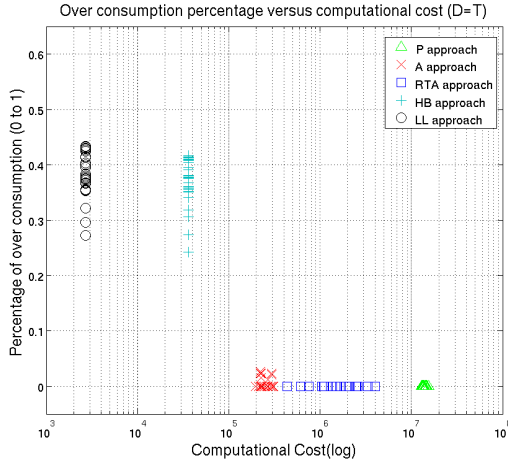


Figure 8: Over-consumption percentage versus computational cost when $D=T$.

The $\mathcal{A}$ approach has a middle cost and a rejection ratio of zero for $D = T$ and for $D \leq T$. Furthermore, the proposed method has predictable costs, and their computational cost is almost 20 times less than the cost of the $\mathcal{P}$ approach. Notice that the gap in the magnitude of the cost between the proposed approach and the $\mathcal{P}$ approach increases with the arrival of tasks.

7.2. Over-consumption percentage versus computational cost

Figures 8 and 9 show the energy over-consumption in function of the computational cost for the worst case. The over-consumptions are due to deviations in the calculation of the scaled factor α compared with the calculation using an exact method. Notice that deviations in alpha lead to quadratic over-consumption – depending on the polynomial expression that describes CPU power consumption.

In Figure 8 we can observe that LL has a small computational cost, but an energy over-consumption close to 45%. HB has an over-consumption slightly lower than LL. When $D \leq T$, the LLM approach shows an over-consumption for the worst case of around 60% to 70%.

Moreover, we obtain an intermediate computational cost when using the reduced approach $\mathcal{A}$. When $D = T$, the proposed method has a higher concentration of points around the 0% for the worst case, and with some values of over-consumption above 0, and which do not exceed 2.5%. When $D \leq T$, we ob-

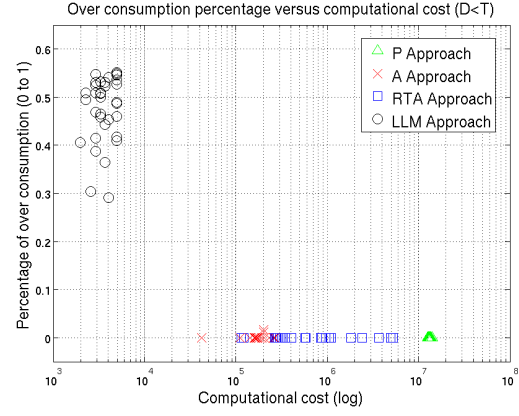


Figure 9: Over-consumption percentage versus computational cost when $D < T$

tain an over-consumption equal to 0 using the $\mathcal{A}$ approach. For both cases, $D = T$ and $D \leq T$, the approaches $\mathcal{P}$ and $\mathcal{RTA}$ show an over-consumption equal to 0, but with higher costs and dispersion respectively.

7.3. Rejection ratio versus over-consumption percentage

Figures 10 and 11 show the percentage of rejected tasks in comparison with energy over-consumption percentages. When

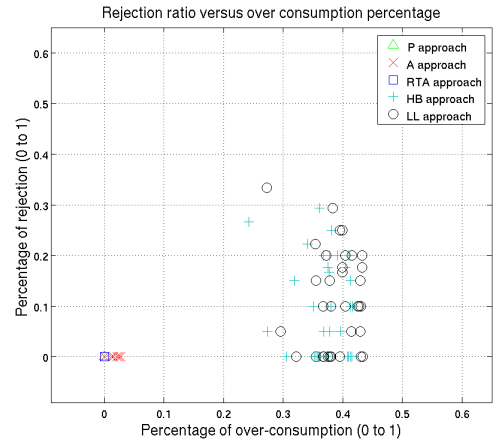


Figure 10: Rejection ratio versus energy over-consumption percentage when $D=T$.

$D = T$, the LL and HB approaches are dispersed at the intersection of the areas of between 20% and 45% of over-consumption and the area of between 0% and 35% of rejection ratio. The increase in the percentage in these approaches is determined by the utilization and the number of tasks. The $\mathcal{A}$ approach has no

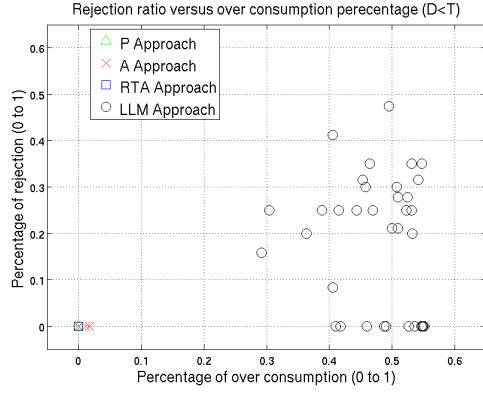


Figure 11: Rejection ratio versus energy over-consumption percentage when $D < T$.

points on the axis of the rejection ratio. On the axis of over-consumption, 5% of the plotted points for $\mathcal{A}$ are between 1.2% and 2.5%. The other points are at 0%. The other approaches have an over-consumption and a rejection ratio equal to 0.

When $D \leq T$, the LLM approach is dispersed at the intersection of the areas between 50% and 70% of over-consumption and the area of between 0% and 70% of the rejection ratio. The $\mathcal{A}$, $\mathcal{P}$ and RTA approaches have an over-consumption and a rejection ratio equal to 0.

7.4. Other simulations

Figure 12 shows separately the evolution of the energy saved with respect to the saving achieved with an exact algorithm in function of the arrival of tasks for each task group (A,B, and C) and their respective arrivals (LL1, LL2 and LL3). It shows a utilization of 95% for the complete tasks set with $D = T$ and $D \leq T$.

Figure 13 shows the evolution of the computational cost of the proposed algorithm when compared with the other static algorithms with respect to the arrival of new tasks for the whole test set when $D = T$. The figure shows the variability of the cost for the RTA approach, which varies depending on the task set and increases with the number tasks arriving.

Figure 13 also shows the evolution of the $\mathcal{P}$ approach, whose cost increases depending on the number of tasks in the order 2^n , where n is the number of tasks. The cost of our static algorithm

is smaller when compared to $\mathcal{P}$ and RTA algorithms, and their evolution with respect to the number of arrivals of tasks is far more muffled than $\mathcal{P}$ and RTA.

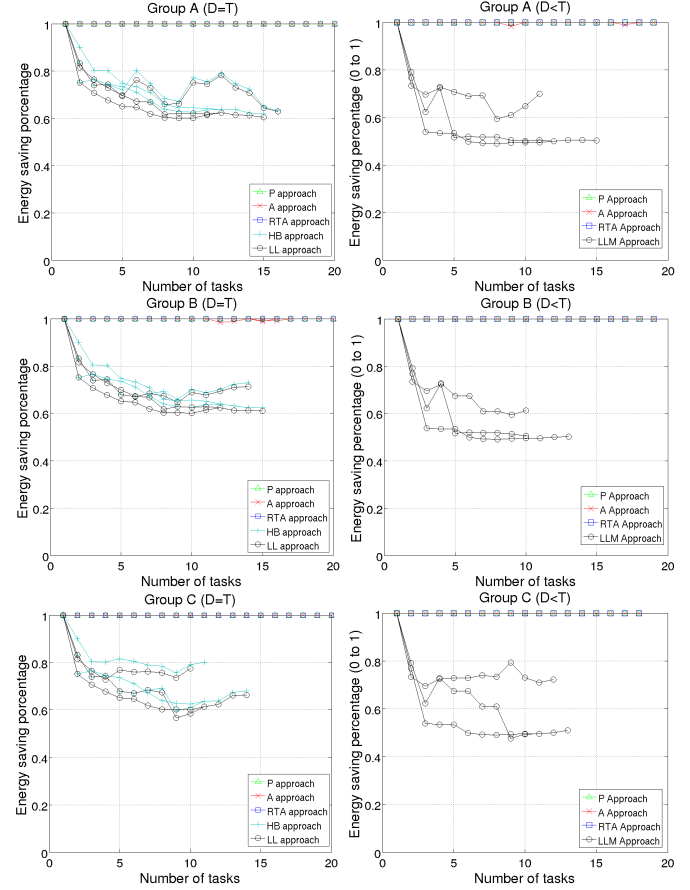


Figure 12: Evolution of the energy saved with respect to the saving with a exact algorithm in function to the arrival of tasks. Utilization equal to 95%

8. Conclusions

In this paper a new algorithm ($\mathcal{A}$ approach) for computing the processor frequency scaling factor under fixed priority assignment with $D = T$ and $D \leq T$ is proposed. This algorithm is performed with a reduced computational cost and minimizes CPU energy consumption while guaranteeing no missed deadlines. Therefore, it can be used as on-line acceptance test in systems with dynamic workload.

Using extensive simulations, we have evaluated the behaviour of several existing feasibility tests that have been adapted to

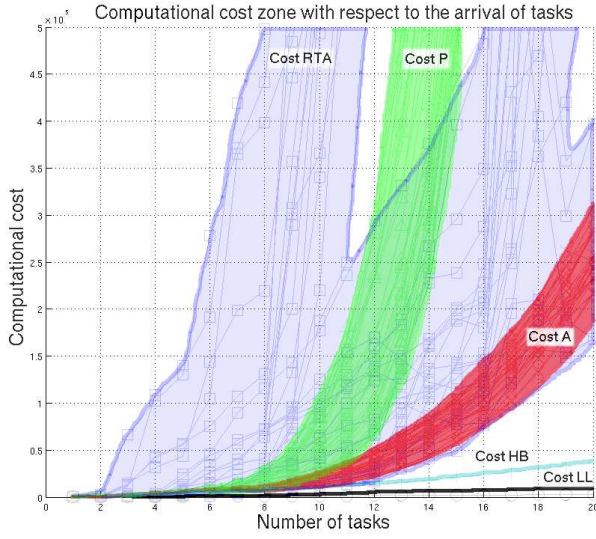


Figure 13: Evolution of the computational cost of the algorithms with respect to the arrival of new tasks.

compute a frequency scaling factor α for the proposed approach $\mathcal{A}$. The analysis has been presented from the point of views of energy consumption, task rejection ratio, and real computing costs.

Our simulation results show that the proposed algorithm outperforms other constant voltage scaling algorithms from the cost, predictability, and task acceptance points of view. However, some small deviations were observed in the energy saved. The proposed approach is compared with other approaches. The proposed algorithm enables not only the high precision verification of system feasibility and the computation of α in an environment with a dynamic processor load. The algorithm can also be used in real-time operating systems because the computational cost represents a small system overload.

The main features of the described approach are presented in Table 2.

As a future work, we plan to investigate the case where memory access and the system bus perform at different clock speeds – together with processor frequency scaling (Marvell, 2008). In this case, the combination of these parameters will enable an even greater reduction in energy consumption for the whole computing system.

Alg.	Cost	Rejection ratio	Over-cons. ratio	D / T
S	Very high	None	None	$D \leq T$
$\mathcal{P}$	High	None	None	$D \leq T$
RTA	Unpredictable Middle-High	None	None	$D \leq T$
$\mathcal{A}^\dagger$	Low-Middle	None	Very low	$D \leq T$
LLM	Low	Very High	Very High	$D \leq T$
HB	Low	High	High	$D = T$
LL	Very Low	High	High	$D = T$

Table 2: Comparison of approaches to compute a constant frequency scaling factor α . († : proposed approach $\mathcal{A}$).

References

- Albertos, P., Crespo, A., & Simó, J. (2006). Control kernel: A key concept in embedded control systems. *4th IFAC Symposium on Mechatronic Systems*, .
- AlEnawy, T., & Aydin, H. (2005). Energy-aware task allocation for rate monotonic scheduling. In *Real Time and Embedded Technology and Applications Symposium, 2005. RTAS 2005. 11th IEEE* (pp. 213–223).
- Anderson, J. H., Bud, V., & Devi, U. C. (2005). An edf-based scheduling algorithm for multiprocessor soft real-time systems. In *ECRTS '05: Proceedings of the 17th Euromicro Conference on Real-Time Systems* (pp. 199–208). Washington, DC, USA: IEEE Computer Society.
- Aydin, H., Devadas, V., & Zhu, D. (2006). System-level energy management for periodic real-time tasks. In *RTSS '06: Proceedings of the 27th IEEE International Real-Time Systems Symposium* (pp. 313–322). Washington, DC, USA: IEEE Computer Society.
- Aydin, H., Melhem, R., Mossé, D., & Mejía-Alvarez, P. (2004). Power-aware scheduling for periodic real-time tasks. *IEEE Trans. Comput.*, 53, 584–600.
- Bini, E., & Buttazzo, G. (2004). Schedulability analysis of periodic fixed priority systems. *IEEE Transactions on Computers*, 53, 1462–1473.
- Bini, E., Buttazzo, G., & Lipari, G. (2005). Speed modulation in energy-aware real-time systems. In *ECRTS '05: Proceedings of the 17th Euromicro Conference on Real-Time Systems* (pp. 3–10). Washington, DC, USA: IEEE Computer Society.
- Bini, E., Buttazzo, G. C., & Buttazzo, G. M. (2003). Rate monotonic analysis: the hyperbolic bound. *IEEE Transactions on Computers*, 52, 933–942.
- Bini, E., Natale, M. D., & Buttazzo, G. C. (2006). Sensitivity analysis for fixed-priority real-time systems. In *Proceedings of the 18th Euromicro Conference on Real-Time Systems*. Dresden, Germany.
- Briao, E. W., Barcelos, D., Wronski, F., & Wagner, F. R. (2007). Impact of task migration in noc-based mpsoes for soft real-time applications. In *Very Large Scale Integration, 2007. VLSI - SoC 2007. IFIP International Conference on* (pp. 296–299).
- Cervin, A. (2003). *Integrated Control and Real-Time Scheduling*. Ph.D. thesis

- Department of Automatic Control, Lund University ISRN LUTFD2/TFRT-1065-SE.
- Chen (2007). Energy-efficient real-time task scheduling with task rejection. .
- Chen, J.-J. (2005). Multiprocessor energy-efficient scheduling for real-time tasks with different power characteristics. In *ICPP '05: Proceedings of the 2005 International Conference on Parallel Processing* (pp. 13–20). Washington, DC, USA: IEEE Computer Society.
- Cho, Y., & Chang, N. (2006). Energy-aware clock-frequency assignment in microprocessors and memory devices for dynamic voltage scaling. *Computer-Aided Design of Integrated Circuits and Systems, IEEE Transactions on*, 26, 1030–1040.
- Crespo, A., Albertos, P., Balbestre, P., Vall'es, M., Lluesma, M., & Sim'o, J. (2006). Schedulability issues in complex embedded control systems. *IEEE International Conference on Control Applications*, .
- Duan, C., & Khatri, S. P. (2006). Computing during supply voltage switching in dvs enabled real-time processors. In *IEEE ISCAS Conference* (p. 2254).
- Emberson, P., & Bate, I. (2007). Minimising task migration and priority changes in mode transitions. In *RTAS '07: Proceedings of the 13th IEEE Real Time and Embedded Technology and Applications Symposium* (pp. 158–167). Washington, DC, USA: IEEE Computer Society.
- Gaujal, B., & Navet, N. (2007). Dynamic voltage scaling under edf revisited. .
- Jejurikar, R., & Gupta, R. (2004). Dynamic voltage scaling for systemwide energy minimization in real-time embedded systems. In *ISLPED '04: Proceedings of the 2004 international symposium on Low power electronics and design* (pp. 78–81). New York, NY, USA: ACM.
- J.L.Posadas, J.L.Poza, J.E.Sim, G.Benet, & F.Blanes (2008). Agent based distributed architecture for mobile robot control. *Engineering Applications of Artificial Intelligence*, 21, 805–823.
- Kim, W., Shin, D., Yun, H.-S., Kim, J., & Min, S. L. (2002). Performance comparison of dynamic voltage scaling algorithms for hard real-time systems. In *RTAS '02: Proceedings of the Eighth IEEE Real-Time and Embedded Technology and Applications Symposium (RTAS'02)* (p. 219). Washington, DC, USA: IEEE Computer Society.
- Kumar, G. S. A., & Manimaran, G. (2007). Energy-aware scheduling of real-time tasks in wireless networked embedded systems. In *RTSS '07: Proceedings of the 28th IEEE International Real-Time Systems Symposium* (pp. 15–24). Washington, DC, USA: IEEE Computer Society.
- Kumar, G. S. A., Manimaran, G., & Wang, Z. (2008). End-to-end energy management in networked real-time embedded systems. *IEEE Trans. Parallel Distrib. Syst.*, 19, 1498–1510.
- Lee, E. A. (2006). Cyber-physical systems – are computing foundations adequate? *NFS Workshop On Cyber-Physical Systems*, .
- Lehoczy, J., Sha, L., & Ding, Y. (1989). The rate-monotonic scheduling algorithm: Exact characterization and average case behavior. In *10th IEEE Real-Time Systems Symposium* (pp. 166–172).
- Leung, L.-F. (2005). Exploiting dynamic workload variation in low energy. .
- Li, T., & Ding, C. (2001). Instruction balance and its relation to program energy consumption. In *Internaional Workshop Languages and Compilers for Parallel computing*.
- Liang, W.-Y., Chen, S.-C., Chang, Y.-L., & Fang, J.-P. (2008). Memory-aware dynamic voltage and frequency prediction for portable devices. In *Embedded and Real-Time Computing Systems and Applications, 2008. RTCSA '08. 14th IEEE International Conference on* (pp. 229–236).
- Liu, Y., & Mok, A. (2003). An integrated approach for applying dynamic voltage scaling to hard real-time systems. In *9th IEEE Real-Time and Embedded Technology and Applications Symposium*. Toronto, Canada.
- Marinoni, M., & Buttazzo, G. (2007). Elastic dvs management in processors with discrete voltage/frequency modes. *IEEE Transactions on Industrial Informatics*, 3.
- Marvell (2008). System and timer configuration developers manual. In *Marvell PXA3xx Processor Family*. volume I.
- Mejia-Alvarez, P., Levner, E., & Mosse, D. (2002). Power-optimized scheduling server for real-time tasks. In *Real-Time and Embedded Technology and Applications Symposium, 2002. Proceedings. Eighth IEEE* (pp. 239–250).
- Mejia-Alvarez, P., Levner, E., & Mossé, D. (2004). Adaptive scheduling server for power-aware real-time tasks. *ACM Trans. Embed. Comput. Syst.*, 3, 284–306.
- Piao, X., Kim, H., Cho, Y., Park, M., Han, S., Park, M., & Cho, S. (2009). Energy consumption optimization of real-time embedded systems. In *ICESS '09: Proceedings of the 2009 International Conference on Embedded Software and Systems* (pp. 281–287). Washington, DC, USA: IEEE Computer Society.
- Pillai, P., & Shin, K. (2001). Real-time dynamic voltage scaling for low-power embedded operating systems. In *ACM symposium on operating systems principles*. Banff, Canada.
- Pouwelse, J., Langendoen, K., & Sips, H. (2001). Dynamic voltage scaling on a low-power microprocessor. In *Seventh Internacional Conference Mobile Computing and Networking (MOBICOM)*.
- Quan (2004). Fixed priority scheduling for reducing overall energy on variable voltage processors. .
- Racu, R., Jersak, M., & Ernst, R. (2005). Applying sensitivity analysis in real-time distributed systems. In *In 11th IEEE Real-Time Technology and Applications Symposium (RTAS'05)* (pp. 160–169). IEEE Computer Society.
- Real, J., & Crespo, A. (2004). Mode change protocols for real-time systems: A survey and a new proposal. *Real-Time Syst.*, 26, 161–197.
- Saewong, S., & Rajkumar, R. R. (2003). Practical voltage-scaling for fixed-priority rt-systems. In *RTAS '03: Proceedings of the The 9th IEEE Real-Time and Embedded Technology and Applications Symposium* (p. 106). Washington, DC, USA: IEEE Computer Society.
- Scordino, C., & Lipari, G. (2006). A resource reservation algorithm for power-aware scheduling of periodic and aperiodic real-time tasks. *IEEE Trans. Comput.*, 55, 1509–1522.
- Sha, L., Abdelzaher, T., Arzn, K.-E., Cervin, A., Baker, T., Burns, A., Buttazzo, G., Caccamo, M., Lehoczy, J., & Mok, A. K. (2004). Real-time scheduling theory: A historical perspective. In *Real-Time Systems* 28:23 (p. 101155).
- Simarro, R., Coronel, J., Simó, J., & Blanes, J. (2008). Hierarchical and dis-

- tributed embedded control kernel. In *17th IFAC World Congress*. Seoul, Korea.
- Sloss, A. N., Symes, D., & Wright, C. (2004). *Arm System Developer's Guide (Designing and Optimizing System Software)*. (Primera ed.). Elsevier.
- Tavares, E., Maciel, P., Silva, B., & Oliveira, M. N., Jr. (2008). Hard real-time tasks' scheduling considering voltage scaling, precedence and exclusion relations. *Inf. Process. Lett.*, 108, 50–59.
- Wolf, W. (2009). Cyber-physical system. *IEEE Computer. Innovative Techonolgy for Computer Professionals.*, 42, 88–89.
- Xia (2008). Control-theoretic dynamic voltage scaling for embedded controllers, .
- Yazdi, H., Salmani, H., Khatib-Astaneh, N., Salmani, M., & Fard, A. (2008). Exploiting laxity for heterogeneous multiprocessor real-time scheduling. In *Information and Communication Technologies: From Theory to Applications, 2008. ICTTA 2008. 3rd International Conference on* (pp. 1–6).
- Yi, J., Poellabauer, C., Hu, X. S., Simmer, J., & Zhang, L. (2009). Energy-conscious co-scheduling of tasks and packets in wireless real-time environments. In *RTAS '09: Proceedings of the 2009 15th IEEE Real-Time and Embedded Technology and Applications Symposium* (pp. 265–274). Washington, DC, USA: IEEE Computer Society.
- Yuan, Z., & Wang, G. (2008). Power management for real-time tasks in wireless networked embedded systems. *Computer Science and Software Engineering, International Conference on*, 4, 118–121.
- Yun, H.-S., & Kim, J. (2003). On energy-optimal voltage scheduling for fixed-priority hard real-time systems. *ACM Trans. Embed. Comput. Syst.*, 2, 393–430.
- Zhong, Xiliang, Xu, & Cheng-Zhong (2007). Frequency-aware energy optimization for real-time periodic and aperiodic tasks. *SIGPLAN Not.*, 42, 21–30.
- Zhu, D., Melhem, R., & Mossé, D. (2004). The effects of energy management on reliability in real-time embedded systems. In *International Conference on Computer-Aided Design*.