



Asymptotic Performance and Energy Consumption of SLACK

Anne Benoit, Louis-Claude Canon, Redouane Elghazi, Pierre-Cyrille Heam

► To cite this version:

Anne Benoit, Louis-Claude Canon, Redouane Elghazi, Pierre-Cyrille Heam. Asymptotic Performance and Energy Consumption of SLACK. RR-9501, Inria Lyon. 2023, pp.27. hal-04021482v2

HAL Id: hal-04021482

<https://inria.hal.science/hal-04021482v2>

Submitted on 7 Jun 2023

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



Distributed under a Creative Commons Attribution 4.0 International License



Asymptotic Performance and Energy Consumption of SLACK

Anne Benoit, Louis-Claude Canon, Redouane Elghazi,
Pierre-Cyrille Héam

**RESEARCH
REPORT**

N° 9501

March 2023

Project-Team ROMA

ISRN INRIA/RR--9501--FR+ENG

ISSN 0249-6399



Asymptotic Performance and Energy Consumption of Slack

Anne Benoit^{*}, Louis-Claude Canon[†], Redouane Elghazi^{†*},
Pierre-Cyrille Héam[†]

Project-Team ROMA

Research Report n° 9501 — version 2 — initial version March 2023 —
revised version May 2023 — [24](#) pages

Abstract: Scheduling n independent tasks onto m identical processors in order to minimize the makespan has been widely studied. As an alternative to classical heuristics, the SLACK algorithm groups tasks by packs of m tasks of similar execution times, and schedules first the packs with the largest differences. It turns out to be very performant in practice, but only few studies have been conducted on its theoretical properties. We derive novel analytical results for SLACK, and in particular, we study the performance of this algorithm from an asymptotical point of view, under the assumption that the execution times of the tasks follow a given probability distribution. The study is building on a comparison of the most heavily loaded machine compared to the least loaded one. Furthermore, we extend the results when the objective is to minimize the energy consumption rather than the makespan, since reducing the energy consumption of the computing centers is an ever-growing concern for economical and ecological reasons. Finally, we perform extensive simulations to empirically assess the performance of the algorithms with both synthetic and realistic execution time distributions.

Key-words: Scheduling, asymptotic performance, energy consumption, empirical study.

^{*} LIP, ENS Lyon, Equipe Inria Roma, France.

[†] FEMTO-ST, U. Franche-Comté, France.

Performance asymptotique et consommation énergétique de Slack

Résumé : Ordonnancer n tâches indépendantes sur m processeurs identiques afin de minimiser le temps total d'exécution est un problème qui a été largement étudié. Comme alternative aux heuristiques classiques, l'algorithme SLACK groupe les tâches en paquets de m tâches avec un temps similaire, et ordonnance d'abord les paquets avec les plus grandes différences de temps. Cette approche est très efficace en pratique, mais peu d'études se sont intéressées à son étude théorique. Nous proposons de nouveaux résultats analytiques pour SLACK. En particulier, nous étudions la performance de l'algorithme d'un point de vue asymptotique, en supposant que le temps d'exécution des tâches suit une distribution de probabilités connue. L'étude se base sur une comparaison entre le processeur le plus chargé et celui qui l'est le moins. De plus, nous étendons ces résultats à une fonction objective différente: la minimisation de la consommation énergétique. En effet, réduire la consommation énergétique des centres de calcul est un des défis majeurs actuels, à la fois d'un point de vue économique, mais aussi et surtout écologique. Finalement, nous effectuons de nombreuses simulations pour étudier de façon empirique la performance des algorithmes, avec des distributions de temps d'exécution synthétiques d'une part, et réalistes d'autre part.

Mots-clés : Ordonnancement, performance asymptotique, consommation énergétique, étude empirique.

Contents

1	Introduction	4
2	Related work	5
3	Framework	6
4	A Bound for SLACK	8
5	Convergence Speed of SLACK	11
5.1	Convergence of the Makespan	12
5.2	Convergence of the Energy Consumption	13
6	Simulations	18
6.1	Experimental Setting	18
6.2	Simulations: Study of δ_j and β_j	18
6.3	Simulations: Energy minimization	20
7	Conclusion	22

1 Introduction

The problem of minimizing the computation time when scheduling n independent tasks on m identical processors is at the basis of scheduling theory, and a building block for solving many more complicated problems, hence it remains very important even though it has already been widely studied. Using Graham's notation [13], this problem is denoted $P||C_{\max}$.

While the problem is NP-complete (equivalent to 2-partition with two processors, or 3-partition when the number of processors m is part of the input), an easy way to get efficient solutions consist in ordering the n tasks according to some criterion, and then perform a list schedule, i.e., schedule the next task of the list on the least loaded processor, hence never leaving a processor idle. A classic ordering is the one of LPT (Longest Processing Time), which orders tasks from the longest to the smallest [14]. This algorithm has proven to have good theoretical and even better practical performance. In particular, its rate of convergence has been studied, and new results were recently established when the distribution of task costs is generated using uniform integer compositions [5].

More recently, the SLACK heuristic was proposed in [7], showing promising empirical performance compared to LPT. Its principle is based on grouping tasks of similar execution times into packs, sorting the resulting packs by non-decreasing similarity (the similarity of a pack denoting the maximum difference of execution times between its tasks), and then scheduling the tasks in the order determined by the packs, following a list schedule (assign the next task to the least loaded processor). The idea is that a single pack cannot bring the imbalance of the processors too high, and the hope is that the packs balance each other. The objective is that the tasks in the last scheduled packs are very close to each other, hence they will not create a large imbalance at the end of the schedule. While this SLACK algorithm benefits from favorable empirical performance, fewer analyses have been conducted on its theoretical properties.

These heuristics were proposed in order to minimize the makespan, i.e., the maximum execution time among the processors. Another core problem consists in minimizing the *energy consumption*, as the energy consumption of current platforms is an ever-growing concern, both for economical and ecological reasons. To optimize the energy consumption, modern processors can run at different speeds, and their power consumption is then the sum of a static part (the cost for a processor to be turned on) and a dynamic part, which is a strictly convex function of the processor speed. More precisely, a processor running at speed s dissipates a power of s^α Watts, where $2 \leq \alpha \leq 3$ [2]. Hence, a higher speed allows executing a task more rapidly, but at the price of a much higher amount of energy consumed. Finding a schedule now consists in deciding on which processor to execute each task and to decide at which speed the task is executed.

Therefore, we revisit this classic problem of scheduling n independent tasks onto m identical processors, with the aim of deriving analytical results for SLACK, when the goal is to minimize the makespan or the energy consumption. We study the performance of SLACK from an asymptotical point of view, under

the assumption that the execution times of the tasks follow a given probability distribution. The study is building on a comparison of the most heavily loaded machine compared to the least loaded one, and hence it provides interesting insights both for the study of the classic makespan objective function, and its translation to the energy consumption. The goal of this paper is therefore to answer two main questions left unresolved in the literature so far: (i) provide a theoretical study to analyze the performance of SLACK, and (ii) consider the energy consumption in the theoretical and empirical analysis of the algorithms. Our main contributions are the following:

- A fundamental bound related to the result of SLACK (Section 4);
- A convergence rate for the makespan of SLACK when using uniform and exponential distributions, by applying the bound of Section 4 (Section 5.1);
- A general result for bounding the energy consumption (agnostic of the algorithm and the task distribution) and its application to SLACK, by applying the bound of Section 4 (Section 5.2);
- Simulations for comparison with the theoretical bounds that were computed for SLACK and LPT (Section 6).

First, Section 2 summarizes the existing contributions related to either the energy minimization problem or LPT and SLACK. Section 3 presents the problems and algorithms (LPT and SLACK). Then, Section 4 presents a useful bound on the result given by SLACK. Section 5 proposes applications of this bound: theoretical asymptotic results related to the minimization of the makespan and the energy with SLACK. In the case of the energy, Section 5.2 also gives a method to derive energy related guarantees for any algorithm bounded similarly to SLACK in Section 4. Section 6 presents the experimental results of the empirical study of LPT and SLACK. Finally, Section 7 concludes.

2 Related work

Lowering the energy consumption of computational tasks has been widely studied in the last decades, be it in the context of High Performance Computing or in other contexts, such as Cloud Computing. Many models have been proposed for the energy consumption of CPUs. For instance, the energy consumption is scaling quadratically with the speed of the CPU in [20], and there is a focus on the online evaluation of the expected idle time. In [21], the only assumption is that the energy consumption is a convex function of the speed of the CPU, and clairvoyant online and offline solutions are proposed to the problem. The heuristics presented in these two articles are then evaluated, either empirically in [20], or with approximation ratios in [21]. In our work, we explore another way of evaluating algorithms, following the remark that with large systems, stochastic asymptotic results should be relevant.

Recent surveys such as [8] and [19] compile various techniques used for energy-efficient computing, including scheduling techniques. These techniques may use either Dynamic Voltage and Frequency Scaling (DVFS), as in [15], where the frequency (and hence the speed) of processors may be chosen, or Dy-

dynamic Power Management (DPM) as in [4]. These studies propose algorithms, but they mainly focus on an empirical evaluation of these algorithms, without theoretical study.

As for scheduling algorithms that have low complexities (and therefore low energy consumption), LPT has been a well known algorithm for decades and is known to provide good theoretical and practical performance while keeping a low time complexity in $O(n \log n)$ [14]. A more recent algorithm, SLACK, also remains with an $O(n \log n)$ time complexity, while providing results that are sometimes better than LPT [7, 5].

There are multiple results about the asymptotic behavior of LPT under different assumptions. Frenk and Rhinnooy Kan [12] and Coffman et al. [6] study the difference between LPT and the optimal solution in the case where the execution times of the tasks follow a probability distribution of cumulative distribution function of the form $F(x) = x^\alpha$, where $0 < \alpha < +\infty$. Loulou [16] and Piersma and Romeijn [17] do not look at specific distributions, but instead they study LPT under the assumption that the execution times are independent and identically distributed random variables. More recently, Benoit et al. [5] studied the asymptotic optimality of SLACK and LPT under the assumption that the execution times are generated using a distribution called the uniform integer composition.

3 Framework

The $P||C_{\max}$ problem is a classic scheduling problem, where n tasks have to be scheduled on m identical machines, with the objective function of makespan minimization, i.e., minimize the execution time of the machine that completes last ($C_{\max}$). There are no constraints on tasks, which can be assigned to any machine in any order. Each task has a number of operations to perform, that we call its work and denote by w_i , and the time to execute the task is usually $t_i = w_i$, assuming that the machine executes one operation per time unit (speed $s = 1$). The problem complexity is well known, and in particular the associated decision problem is NP-complete as soon as $m \geq 2$.

List scheduling and LPT. In order to solve this $P||C_{\max}$ problem, a simple but effective heuristic algorithm consists in never letting a machine idle, i.e., as soon as a task completes on a machine, a new task is assigned to this machine. This is called *list scheduling*, and it can be implemented as in Algorithm 1, by keeping the load of each machine in a vector $\vec{W}$ of length m initialized to $(0, 0, \dots, 0)$. For each task, we assign it to the currently least loaded machine, and the makespan is the maximum value of the vector $\vec{W}$ at the end of the execution. Any list schedule (whatever the order of tasks) is known to be a $(2 - \frac{1}{m})$ -approximation algorithm [14]. A variant of the List Scheduling heuristic consists in first sorting the list L by non-increasing task works, and it is called *Longest-Processing-Time-first* (LPT for short). This can be used if all tasks are known beforehand (offline scheduling), and it improves the approximation ratio

of the algorithm to $(\frac{4}{3} - \frac{1}{3m})$ [14].

Slack. In this paper, we mainly focus on the SLACK algorithm, that was introduced in [7] and consists in applying the List Scheduling heuristic with a particular pretreatment on the list of tasks, as detailed in Algorithm 2. We first fill the list L to have a number of elements r that is a multiple of m , by adding dummy tasks of work 0. Then, tasks are sorted by non-increasing works and grouped by packs of m tasks, and then the packs are themselves sorted by non-increasing difference between the work of the longest task of the pack and the smallest one (α_i 's). These differences are denoted β_k , where $\beta_1 \geq \beta_2 \dots \geq \beta_{r/m}$. They correspond to the sorted α_i 's.

Let us denote by $c_i(j)$ the load of processor j after $i \times m$ tasks (i.e., the i first packs) have been scheduled. Hence, $c_i(j) = \vec{W}[j]$ after $i \times m$ steps of the loop line 2 of Algorithm 1. One has for instance $c_0(j) = 0$ for all j (initial load), and then at each iteration i , we schedule one more pack with m tasks. We then define $\delta_i = \max_{0 \leq j, j' < m} (|c_i(j) - c_i(j')|)$, which is the maximum difference of load between two processors after iteration i .

Note that these values β_i and δ_i can be extended to any algorithm, in particular LPT, by simply grouping tasks by *packs* of m tasks (in the order in which they are scheduled), and then checking differences between tasks of a same pack (β_i), and differences in the load of processors at each iteration, hence, after m tasks have been scheduled (δ_i).

From makespan to energy consumption. When the goal is to minimize the energy consumption, we further consider that the frequency of the processors can be scaled using DVFS (Dynamic Voltage and Frequency Scaling). Hence, these processors have a static power P_{stat} , and can be operated at any speed (or frequency) $s \in \mathbb{R}_+^*$ [3], while we assumed so far that $s = 1$.

The execution time of task T_i at speed s then becomes $t_{i,s} = \frac{w_i}{s}$. In terms of energy consumption, there is a static part, which corresponds to the power consumed when the m processors are turned on, during a time C_{max} , hence a total of $m \times C_{max} \times P_{stat}$. For each task T_i , there is also a dynamic energy consumption, directly related to the speed s at which the processor operates the task. Using a general model, the dynamic energy consumption is $t_{i,s} \times s^\alpha$ [2], where $\alpha > 1$ (in general, $2 \leq \alpha \leq 3$). Finally, the total energy consumption of

Algorithm 1 ListScheduling(L, m)

Require: List L of n positive floats (task works); Number of processors m .

- 1: Let $\vec{W}$ be a vector of length m initialized to $\vec{W} = (0, 0, \dots, 0)$;
 - 2: **for** $w \in L$ in the order they appear in the list **do**
 - 3: Let j be the index of a minimal element of $\vec{W}$;
 - 4: $\vec{W}[j] = \vec{W}[j] + w$;
 - 5: **end for**
 - 6: **return** $\vec{W}$;
-

Algorithm 2 SLACK (L, m)**Require:** List L of n positive floats (task works); Number of processors $m \leq n$.

- 1: Add $(-n \bmod m)$ elements of work 0 at the end of L ;
- 2: $r = n + (-n \bmod m)$;
- 3: $L' = [x_1, \dots, x_r]$ is obtained by sorting L non-increasingly;
- 4: **for** $0 \leq i \leq \frac{r}{m} - 1$ **do**
- 5: $K_i = [x_{im+1}, x_{im+2}, \dots, x_{im+m}]$;
- 6: $\alpha_i = x_{im+1} - x_{im+m}$;
- 7: **end for**
- 8: Let $H = [\alpha_{i_1}, \dots, \alpha_{i_{\frac{r}{m}}}] = [\beta_1, \dots, \beta_{\frac{r}{m}}]$ be a non-increasing sequencing of the α_i 's; L_{SLACK} is obtained by concatenating the K_i 's in the same order as the α 's in H .
- 9: $\tilde{W} = \text{ListScheduling}(L_{\text{SLACK}}, m)$;

a schedule of length $C_{\max}$, where T_i is operated at speed s_i , is:

$$E = m \times C_{\max} \times P_{\text{stat}} + \sum_{i=1}^n t_{i,s_i} \times s_i^\alpha.$$

For convenience, the main notations are summarized in Table 1.

4 A Bound for SLACK

This section is dedicated to proving a fundamental bound related to the algorithm SLACK.

Let $\mathcal{X}$ be a distribution with positive values. We denote by $C(n, m, \mathcal{X})$ the random variable of the makespan returned by the SLACK algorithm on m processors on a list of n tasks that are independent random variables of distribution $\mathcal{X}$. Let $X_1, \dots, X_n$ be n independent random variables distributed according to $\mathcal{X}$. Let $X_{1:n} \leq X_{2:n} \leq \dots \leq X_{n:n}$ be associated order statistics. Particularly $X_{1:n}$ is the minimum of the X_i 's and $X_{n:n}$ the maximum. Let $D_i = (X_{i:n} - X_{i-1:n})$ for every $1 \leq i \leq n$, with the convention $X_{0:n} = 0$. The D_i 's are classically called spacings of adjacent order statistics. Let $\Delta_{\mathcal{X},n}$ be the random variable of the maximal value of D_i 's, that is the maximal difference between two consecutive $X_{i:n}$ (and between 0 and $X_{1:n}$).

Theorem 1. *When using SLACK (Algorithm 2), $\max_{1 \leq i, j \leq m} (W_i - W_j) \leq m \Delta_{\mathcal{X},n}$.*

Lemma 1. *If at step $j+1$, we put two or more tasks on a processor, then each of these tasks has an execution time at most δ_j .*

Proof. Assume there is a processor a that receives two or more tasks at step $j+1$. It means that the processor b that initially had the highest execution time does not receive any task, otherwise it would mean that we have distributed exactly one task to each processor. We then get that before receiving its second task, the

Symbol	Definition
m	number of processor
n	number of tasks
$\{T_1, \dots, T_n\}$	the n tasks
w_i	work of T_i (corresponding to the number of operations required by the task)
$t_i = w_i$	the execution time of T_i at speed 1
$t_{i,s} = \frac{w_i}{s}$	execution time of T_i at speed s
$m \times C_{\max} \times P_{\text{stat}}$	static energy consumption for a duration $C_{\max}$
$t_{i,s} \times s^\alpha$	dynamic energy consumption of T_i at speed s
δ_i	largest difference between the total execution times of two processors after having processed $i \times m$ tasks (first i packs)
β_i	largest difference between the execution time of any two tasks in pack i
$c_i(j)$	total execution time of processor j after $i \times m$ tasks have been scheduled (first i packs)
$W_j = \sum_{\text{alloc}(i)=j} w_i$	total work (number of operations) on processor j ; $\text{alloc}(i)$ is the processor on which T_i is allocated
$W_{\max} = \max_{1 \leq j \leq m} W_j$	maximal number of operations allocated to a processor
$W = \sum_{1 \leq i \leq n} w_i$	total number of operations to perform
$\vec{W} = (W_1, \dots, W_m)$	the $\vec{\cdot}$ notation is used for m -length vectors (not only for W)
$\ \vec{x}\ _\alpha = \sqrt[\alpha]{\sum x_i^\alpha}$	classic α -norm of a vector

Table 1: Main Notations

processor a has a lower current execution time than processor b , which means that the first task allocated to processor a at this step has an execution time at most the initial difference between these two processors. So, this execution time is at most δ_j , and since the tasks are sorted in a step, all the subsequent tasks added to the processor a this step also fulfill this condition. $\square$

Lemma 2. *For any processors a and b that both receive at least one task at step $j + 1$, we have $|c_j(a) + \tau_{a,1} - (c_j(b) + \tau_{b,1})| \leq \max(\delta_j, \beta_j)$, where $\tau_{a,1}$ (resp. $\tau_{b,1}$) is the execution time of the first task received by a (resp. by b).*

Proof. Assume without loss of generality that processor b receives its first task before processor a . Then, we have $c_j(b) \leq c_j(a)$ and $\tau_{a,1} \leq \tau_{b,1}$. Let $\tau_{\text{diff}} = \tau_{b,1} - \tau_{a,1}$. We have:

$$\begin{aligned}
|c_j(a) + \tau_{a,1} - (c_j(b) + \tau_{b,1})| &= |(c_j(a) - c_j(b)) - \tau_{\text{diff}}| \\
&\leq \max((c_j(a) - c_j(b)) - \tau_{\text{diff}}, \tau_{\text{diff}} - (c_j(a) - c_j(b))).
\end{aligned}$$

Since $0 \leq \tau_{\text{diff}}$ and $0 \leq (c_j(a) - c_j(b))$ we have

$$\begin{aligned} |(c_j(a) + \tau_{a,1}) - (c_j(b) + \tau_{b,1})| &\leq \max((c_j(a) - c_j(b)), \tau_{\text{diff}}) \\ &\leq \max(\delta_j, \beta_j). \end{aligned}$$

□

Lemma 3. *For every j , $\delta_{j+1} \leq \max(\beta_{j+1}, \delta_j)$.*

Proof. Let $p_{\min}$ be the processor minimizing c_{j+1} and $p_{\max}$ be the processor maximizing c_{j+1} . Let $\hat{\tau}_m = c_j(p_{\min})$ and $\hat{\tau}_M = c_j(p_{\max})$ be their respective execution times before adding the new tasks, and $\{\tau_{m,k}\}_{k \leq K_m}$ and $\{\tau_{M,k}\}_{k \leq K_M}$ are the execution times of the tasks allocated to these processors during this step (sorted by non increasing execution times).

By definition of $p_{\min}$ and $p_{\max}$, $\delta_{j+1} = c_{j+1}(p_{\max}) - c_{j+1}(p_{\min}) = (\hat{\tau}_M + \sum \tau_{M,k}) - (\hat{\tau}_m + \sum \tau_{m,k})$. We see different cases depending on the number of tasks on $p_{\max}$, denoted by K_M .

- Case 1: $K_M = 0$. In this case, the maximum processor has not changed and has not received any task this step, so the difference between the processor $p_{\max}$ and the other processors has only reduced at this step. So, we get $\delta_{j+1} \leq \delta_j \leq \max(\beta_{j+1}, \delta_j)$.
- Case 2: $K_M = 1$. In this case, we have:

$$\begin{aligned} \delta_{j+1} &= (\hat{\tau}_M + \tau_{M,1}) - (\hat{\tau}_m + \sum \tau_{m,k}) \\ &\leq (\hat{\tau}_M + \tau_{M,1}) - (\hat{\tau}_m + \tau_{m,1}) && t_{m,k} \geq 0 \\ &\leq |(\hat{\tau}_M + \tau_{M,1}) - (\hat{\tau}_m + \tau_{m,1})| && \text{the quantity is already positive} \\ &\leq \max(\beta_{j+1}, \delta_j) && \text{Lemma 2} \end{aligned}$$

- Case 3: $K_M > 1$. By using Lemma 1, the last task added to processor $p_{\max}$ has size at most δ_j . The last task of $p_{\max}$ would have been added to any processor with lower execution time if possible, so at the end of step $j + 1$, any processor has a total execution time at least:

$$\begin{aligned} T &\geq (\hat{\tau}_M + \sum \tau_{M,k}) - \tau_{M,K_M} \\ &\geq (\hat{\tau}_M + \sum \tau_{M,k}) - \delta_j && \tau_{M,K_M} \leq \delta_j \end{aligned}$$

So we get:

$$\begin{aligned} \delta_{j+1} &= (\hat{\tau}_M + \sum \tau_{M,k}) - (\hat{\tau}_m + \sum \tau_{m,k}) \\ &\leq (\hat{\tau}_M + \sum \tau_{M,k}) - ((\hat{\tau}_m + \sum \tau_{m,k}) - \delta_j) \\ &\leq \delta_j \\ &\leq \max(\beta_{j+1}, \delta_j) \end{aligned}$$

So, in all cases, we get $\delta_{j+1} \leq \max(\beta_{j+1}, \delta_j)$, proving the lemma. $\square$

Lemma 4. *When using SLACK (Algorithm 2), for every j , $\delta_j \leq \beta_1$.*

Proof. The proof is done by induction on j . For the first iteration ($j = 1$), since the m first considered tasks are each assigned to a different processor, we have:

$$\begin{aligned} \delta_1 &= \max_{p, p'} (|c_1(p) - c_1(p')|) \\ &= \max_{k, k' \in [1, m]} (|x_{k+i_1 \times r/m} - x_{k'+i_1 \times r/m}|) \quad \text{replacing with the first tasks} \\ &= \alpha_{i_1} = \beta_1. \end{aligned}$$

Assume now that $\delta_j \leq \beta_1$ for a $j < r/m$. Using Lemma 3, we can bound $\delta_{j+1} \leq \max(\beta_{j+1}, \delta_j)$. Therefore, $\delta_{j+1} \leq \max(\beta_{j+1}, \beta_1)$. Since $\beta_1 \geq \beta_{j+1}$, we have $\delta_{j+1} \leq \beta_1$, which concludes the proof. $\square$

We are now ready to prove Theorem 1.

Proof. We prove an upper bound for β_1 .

$$\begin{aligned} \beta_1 &= \alpha_{i_1} \\ &= \max_{1 \leq i \leq m} (x_{1+i \times r/m} - x_{m+i \times r/m}) \\ &= \max_{1 \leq i \leq m} \left(x_{1+i \times r/m} - x_{m+i \times r/m} - \sum_{i=2}^{m-1} x_{1+i \times r/m} + \sum x_{1+i \times r/m} \right) \\ &= \max_{1 \leq i \leq m} \left(\sum_{i=1}^{m-1} x_{1+i \times r/m} - x_{1+(i+1) \times r/m} \right) \\ &\leq (m-1) \times \max_{1 \leq i \leq r-1} X_{i+1:n} - X_{i:n} \\ &\leq (m-1) \times \Delta_{\mathcal{X},n} \end{aligned}$$

We just proved that

$$\beta_1 \leq (m-1) \times \Delta_{\mathcal{X},n}. \quad (1)$$

We conclude the proof by combining Equation (1) and Lemma 4 (since $\delta_n = \max_{1 \leq i, j \leq m} (W_i - W_j)$). $\square$

5 Convergence Speed of SLACK

In this section, we use the fundamental bound found in Section 4 to derive asymptotic results on the optimality of SLACK, first in terms of makespan in Section 5.1, and then in terms of energy consumption in Section 5.2.

5.1 Convergence of the Makespan

This section is dedicated to prove asymptotic results on the optimality of SLACK. The following main result is a direct application of Theorem 1:

Proposition 1. *The makespan of SLACK differs from the optimal one by at most $m\Delta_{\mathcal{X},n}$:*

$$0 \leq C(n, m, \mathcal{X}) - OPT \leq \frac{(m-1)^2}{m} \Delta_{\mathcal{X},n} \leq m\Delta_{\mathcal{X},n}.$$

Proof. Since $OPT \geq \frac{1}{m} \sum_{i=1}^r x_i$, one has

$$\begin{aligned} m \times C(n, m, \mathcal{X}) &\leq (m-1) \times \delta_n + \sum_{i=1}^r x_i \\ &\leq (m-1) \times \delta_n + m \times OPT. \end{aligned}$$

Consequently (by Theorem 1),

$$0 \leq C(n, m, \mathcal{X}) - OPT \leq \frac{(m-1)^2}{m} \Delta_{\mathcal{X},n} \leq m\Delta_{\mathcal{X},n}.$$

□

It is worth noting that for many bounded distributions, Proposition 1 will provide results on the convergence of the absolute error $(C(n, m, \mathcal{X}) - OPT)$ when n goes to infinity, as $\Delta_{\mathcal{X},n}$ is smaller when that execution times of the tasks gets denser.

Now, we will use known results on order spacings to obtain convergence results for SLACK. It is proved in [18, corollary 1.4], [1, Section 3] that

$$\mathbb{E}(\Delta_{\mathcal{U}[0,1],n}) \sim \frac{\ln n}{n+1}, \quad (2)$$

where $\mathcal{U}[0,1]$ is the uniform distribution between 0 and 1.

From Proposition 1 and Equation (2), one has the following result, proving that for a fixed m , the SLACK algorithm provides a scheduling that converges in expectation to the optimal (for the makespan):

Corollary 1. *For any fixed $m \geq 2$,*

$$0 \leq \mathbb{E}(C(n, m, \mathcal{U}[0,1])) - \mathbb{E}(OPT) = O\left(m \frac{\ln n}{n+1}\right).$$

A similar result can be obtained for the exponential distribution. It is shown in [10] that, almost surely,

$$\limsup_{n \rightarrow +\infty} \left(\frac{\Delta_{\mathcal{E}_1,n}}{\ln \ln n} \right) = 1, \quad (3)$$

where $\mathcal{E}_1$ is the exponential distribution (with rate 1).

Using Proposition 1 and Equation (3), we then have the following result:

Corollary 2. *For any fixed $m \geq 2$, one has almost surely*

$$0 \leq \limsup_{n \rightarrow +\infty} \left(\frac{C(n, m, \mathcal{E}_1) - OPT}{m \ln \ln n} \right) \leq 1.$$

Corollary 2 does not show a convergence of the makespan of SLACK to the optimal, but that, almost surely, the gap between their difference is under control since $\ln \ln n$ has a very slow growing speed.

5.2 Convergence of the Energy Consumption

Building upon the previous results bounding the δ_i 's for SLACK and analyzing its impact on the makespan, we now move to the problem of minimizing the total energy consumption E , where the speed of each processor can take any value in $\mathbb{R}_+^*$. The main result, stated in Theorem 2, shows how to adapt a classic scheduling algorithm (without speed and energy consideration) into an energy-oriented one. The quality of the solution is bounded by a factor depending on the maximal difference δ between the execution times of the last finishing processor and the first finishing processor.

We start with a preliminary lemma that further defines the shape of optimal solutions: each processor has a constant speed and all processors finish at the same time.

Lemma 5. *In an optimal solution, each processor has a constant speed, and all processors finish at the same time.*

Proof. We first prove that each processor has a constant speed, and then that all processors finish at the same time.

Each processor has a constant speed. Let us assume that there is a processor that does not have constant speed. It means that there are two consecutive amount of work w_1, w_2 being processed at different speeds s_1, s_2 . Let $\gamma_1 = \frac{1}{s_1}$, $\gamma_2 = \frac{1}{s_2}$, $\gamma = \frac{w_1 \times \gamma_1 + w_2 \times \gamma_2}{w_1 + w_2}$. Notice that γ is the weighted average of γ_1 and γ_2 (i.e., $\gamma = t \times \gamma_1 + (1 - t) \times \gamma_2$ with $t = \frac{w_1}{w_1 + w_2}$). By running both amounts of works at speed $\frac{1}{\gamma}$, the total execution time does not change, so the static energy does not change. By strict convexity of $f(x) \mapsto \gamma^{1-\alpha}$ applied to γ_1, γ_2 , and γ , we get that the dynamic energy decreases. So, the total energy consumption decreases.

All processors finish at the same time. Let us assume that the processors do not finish at the same time. Then, there is a processor finishing before at least one other processor. We apply a factor $c < 1$ to all the speeds on this processor such that the $C_{\max}$ does not increase. The static energy does not change, but the dynamic energy decreases. So the total energy consumption decreases. $\square$

Theorem 2. *If an algorithm without speeds outputs a schedule with $\max(W_i - W_j) = \delta$, then we can transform it in polynomial time, with the optimal choice of speeds, into a schedule with $E \leq (1 + \frac{m\delta}{W})\text{OPT}$, where OPT is the minimal energy consumption that could be attained.*

Proof. Using Lemma 5, if the assignment of tasks to processors is fixed, then we only have to choose a constant σ such that processor j runs at speed $\frac{\sigma \times W_j}{W_{\max}}$. The energy we get for a given σ is then:

$$\begin{aligned} E_{W_1, \dots, W_m}(\sigma) &= \frac{W_{\max}}{\sigma} \times m \times P_{\text{stat}} + \sum_j W_j \times \frac{W_{\max}}{\sigma \times W_j} \times \left(\frac{\sigma \times W_j}{W_{\max}} \right)^\alpha \\ &= \frac{1}{\sigma} \times W_{\max} \times m \times P_{\text{stat}} + \sigma^{\alpha-1} \times \sum_j \frac{W_j^\alpha}{W_{\max}^{\alpha-1}}. \end{aligned}$$

This energy consumption is minimized for $\sigma = \frac{W_{\max} \sqrt[\alpha]{m \times P_{\text{stat}}}}{\|\vec{W}\|_\alpha^{\frac{\alpha}{\alpha-1}}}$.

Now, the minimal energy $E_{W_1, \dots, W_m}^{(\min)}$ for this task assignment is:

$$E_{W_1, \dots, W_m}^{(\min)} = (m \times P_{\text{stat}})^{\frac{\alpha-1}{\alpha}} \times \left((\alpha-1)^{\frac{1}{\alpha}} + (\alpha-1)^{\frac{1-\alpha}{\alpha}} \right) \times \|\vec{W}\|_\alpha.$$

We now prove that over all valid $\vec{W}$, this quantity is bounded by $E_{\frac{W}{m}, \dots, \frac{W}{m}}^{(\min)}$. We do so through induction over m .

We have:

$$\left\| \frac{W}{m}, \frac{W}{m}, \dots, \frac{W}{m} \right\|_\alpha = \sqrt[\alpha]{m} \frac{W}{m}.$$

Let $f(W, m) = \min_{\vec{W} \in \mathbb{R}^m | \sum W_i = W} \|\vec{W}\|_\alpha$.

For $m = 1$, we have:

$$f(W, 1) = W = \sqrt[\alpha]{m} \frac{W}{m}.$$

Now, we assume that for $m \geq 1$, we have $f(W, m) = \sqrt[\alpha]{m} \frac{W}{m}$.

By definition of f ,

$$f(W, m+1) = \min_{\vec{W} \in \mathbb{R}^{m+1} | \sum_{i=1}^{m+1} W_i = W} \|\vec{W}\|_\alpha.$$

Therefore, fixing W_{m+1} and by definition of the α -norm,

$$\begin{aligned}
f(W, m+1) &= \min_{w \in [0, W]} \min_{\vec{W} \in \mathbb{R}^{m+1} \mid \sum_{i=1}^{m+1} W_i = W-w \text{ and } W_{m+1}=w} \|\vec{W}\|_\alpha \\
&= \min_{w \in [0, W]} \min_{\vec{W} \in \mathbb{R}^m \mid \sum_{i=1}^m W_i = W-w} \sqrt[\alpha]{w^\alpha + \sum_{i=1}^m W_i^\alpha} \\
&= \sqrt[\alpha]{\min_{w \in [0, W]} w^\alpha + \min_{\vec{W} \in \mathbb{R}^m \mid \sum_{i=1}^m W_i = W-w} \sum_{i=1}^m W_i^\alpha}.
\end{aligned}$$

Now, by definition of f and using the induction hypotheses, we obtain:

$$\begin{aligned}
f(W, m+1) &= \sqrt[\alpha]{\min_{w \in [0, W]} w^\alpha + f(W-w, m)^\alpha} \\
&= \sqrt[\alpha]{\min_{w \in [0, W]} w^\alpha + m \frac{(W-w)^\alpha}{m^\alpha}} \\
&= \frac{\sqrt[\alpha]{m}}{m} \sqrt[\alpha]{\min_{w \in [0, W]} m^{\alpha-1} w^\alpha + (W-w)^\alpha}.
\end{aligned}$$

We now need to study the variations of $g(w) = m^{\alpha-1} w^\alpha + (W-w)^\alpha$ to find the minimum value of $f(W, m+1)$. We have:

$$g'(w) = \alpha \times m^{\alpha-1} \times w^{\alpha-1} - \alpha \times (W-w)^{\alpha-1},$$

so we have:

$$\begin{aligned}
g'(w) < 0 &\Leftrightarrow \alpha \times m^{\alpha-1} \times w^{\alpha-1} - \alpha \times (W-w)^{\alpha-1} < 0 \\
&\Leftrightarrow \alpha \times m^{\alpha-1} \times w^{\alpha-1} < \alpha \times (W-w)^{\alpha-1}.
\end{aligned}$$

Therefore, since $\alpha > 0$, $g'(w) < 0 \Leftrightarrow m^{\alpha-1} \times w^{\alpha-1} < (W-w)^{\alpha-1}$. Now, using that $x \mapsto x^{\alpha-1}$ is an increasing function,

$$g'(w) < 0 \Leftrightarrow m \times w < W - w \Leftrightarrow w < \frac{W}{m+1}.$$

Meaning that the minimum value for $g(w)$ is reached for $w = \frac{W}{m+1}$, so we have:

$$f(W, m+1) = \sqrt[\alpha]{m+1} \frac{W}{m+1}.$$

So we have that $E_{\frac{W}{m}, \dots, \frac{W}{m}}^{(min)}$ is a lower bound of the energy consumption of a schedule, with:

$$E_{\frac{W}{m}, \dots, \frac{W}{m}}^{(min)} = (m \times P_{stat})^{\frac{\alpha-1}{\alpha}} \times \left((\alpha-1)^{\frac{1}{\alpha}} + (\alpha-1)^{\frac{1-\alpha}{\alpha}} \right) \times \sqrt[\alpha]{m+1} \frac{W}{m+1}.$$

In particular, as there exists a schedule with energy consumption OPT, we have:

$$\text{OPT} = E_{W_1^{\text{OPT}}, \dots, W_m^{\text{OPT}}}^{(\min)}$$

that is minored by $E_{\frac{W}{m}, \dots, \frac{W}{m}}^{(\min)}$, meaning that:

$$P_{\text{stat}}^{\frac{\alpha-1}{\alpha}} \times \left[(\alpha-1)^{\frac{1}{\alpha}} + (\alpha-1)^{\frac{1-\alpha}{\alpha}} \right] \times W \leq \text{OPT}.$$

As for the worst case, we know that $\min W_j \leq \frac{W}{m}$. Let $\overrightarrow{W_{\text{worst}}}$ be the vector maximizing the α -norm under the constraint that $W_j - W_i \leq \delta$. For this vector, we have $\max W_j \leq \frac{W}{m} + \delta$, so we get:

$$\left\| \overrightarrow{W_{\text{worst}}} \right\| \leq \sqrt[\alpha]{m} \left(\frac{W}{m} + \delta \right).$$

Now, if an algorithm $\mathcal{A}$ produces a schedule with energy $E_{\mathcal{A}}$ such that $W_j - W_i \leq \delta$ for all i, j , then we have:

$$\begin{aligned} \frac{E_{\mathcal{A}}}{E_{\text{OPT}}} &\leq \frac{\frac{W}{m} + \delta}{\frac{W}{m}} \\ &\leq \frac{W + m \times \delta}{W} \\ &\leq 1 + m \times \frac{\delta}{W}, \end{aligned}$$

which concludes the proof. $\square$

Proposition 2. *The energy consumption of SLACK differs from the optimal one by at most $m^2 \Delta_{\mathcal{X}, n} \frac{\text{OPT}}{W}$:*

$$0 \leq \frac{E(n, m, \mathcal{X}) - \text{OPT}}{\text{OPT}} \leq \frac{m^2 \Delta_{\mathcal{X}, n}}{W}.$$

Proof. Theorem 1 shows that for $\delta = (m-1) \times \Delta_{\mathcal{X}, n}$, we have $\max_{1 \leq i, j \leq n} (W_i - W_j) \leq \delta$. Now using Theorem 2 with this premise, we directly get the desired result. $\square$

Analogously to Proposition 1, Proposition 2 provides asymptotic results on SLACK used for optimizing the energy consumption. Further results can be obtained both for the uniform distribution in Corollary 3 and for the exponential distribution in Corollary 4. Intuitively, the result shows that the relative difference between the energy provided by the adapted SLACK algorithm and the optimal energy consumption converges to 0 almost surely, when $n \rightarrow +\infty$, with a speed at least $\frac{m^2 \log n}{n^2}$ for the uniform distribution and $\frac{m^2 \log \log n}{n}$ for an exponential distribution.

It is proved in [9] that, almost surely,

$$\limsup_{n \rightarrow +\infty} \left(\frac{n \Delta_{\mathcal{U}[0,1], n} - \ln n}{2 \log n} \right) = 1.$$

Corollary 3. *When using SLACK as a base scheduling algorithm with the speed strategy exposed in Lemma 5 with uniform distribution for the tasks, one has almost surely*

$$\limsup_{n \rightarrow +\infty} \left(\frac{E_{\text{SLACK}}(n, m, \mathcal{U}[0, 1]) - \text{OPT}}{\text{OPT}} \times \frac{n^2}{2(2 + \ln 2)m^2 \log n} \right) \leq 1.$$

Proof. First, we simply rewrite the result from [9]:

$$\limsup_{n \rightarrow +\infty} \left(\frac{n\Delta_{\mathcal{U}[0,1],n} - \ln n}{2 \log n} \right) = 1 \quad (4)$$

$$\limsup_{n \rightarrow +\infty} \left(\frac{n\Delta_{\mathcal{U}[0,1],n}}{2 \log n} - \frac{\ln n}{2 \log n} \right) = 1 \quad (5)$$

$$\limsup_{n \rightarrow +\infty} \left(\frac{n\Delta_{\mathcal{U}[0,1],n}}{2 \log n} - \frac{\ln 2}{2} \right) = 1 \quad (6)$$

$$\limsup_{n \rightarrow +\infty} \left(\frac{n\Delta_{\mathcal{U}[0,1],n}}{2 \log n} \right) = 1 + \frac{\ln 2}{2} \quad (7)$$

$$\limsup_{n \rightarrow +\infty} \left(\frac{n\Delta_{\mathcal{U}[0,1],n}}{(2 + \ln 2) \log n} \right) = 1 \quad (8)$$

Using Proposition 2, we have

$$\frac{E_{\text{SLACK}}(n, m, \mathcal{U}[0, 1]) - \text{OPT}}{\text{OPT}} \leq \frac{m^2 \Delta_{\mathcal{U}[0,1],n}}{W} \quad (9)$$

$$\frac{E_{\text{SLACK}}(n, m, \mathcal{U}[0, 1]) - \text{OPT}}{\text{OPT}} \times \frac{n^2}{2(2 + \ln 2)m^2 \log n} \leq \frac{n}{2W} \times \frac{n\Delta_{\mathcal{U}[0,1],n}}{(2 + \ln 2) \log n} \quad (10)$$

Now using Equation 8, and the fact that $\lim_{n \rightarrow +\infty} \frac{n}{2W} = 1$ almost surely with the law of large numbers, we get that almost surely

$$\limsup_{n \rightarrow +\infty} \left(\frac{n}{2W} \times \frac{n\Delta_{\mathcal{U}[0,1],n}}{(2 + \ln 2) \log n} \right) = 1 \quad (11)$$

So using Equation (10), we get that almost surely

$$\limsup_{n \rightarrow +\infty} \left(\frac{E_{\text{SLACK}}(n, m, \mathcal{U}[0, 1]) - \text{OPT}}{\text{OPT}} \times \frac{n^2}{2(2 + \ln 2)m^2 \log n} \right) \leq 1. \quad (12)$$

□

It is proved in [10] that if $\mathcal{E}_1$ is the exponential distribution of rate 1, then, almost surely,

$$\limsup_{n \rightarrow +\infty} \left(\frac{\Delta_{\mathcal{E}_1,n}}{\ln \ln n} \right) = 1.$$

As the rate λ of an exponential distribution is a scaling parameter, we get that if $\mathcal{E}_\lambda$ is the exponential distribution of rate λ , then almost surely

$$\limsup_{n \rightarrow +\infty} \left(\frac{\lambda \Delta_{\mathcal{E}_\lambda,n}}{\ln \ln n} \right) = 1.$$

Corollary 4. *When using SLACK as a base scheduling algorithm with the speed strategy exposed in Lemma 5 with exponential distribution of rate λ for the tasks, for any fixed $m \geq 2$, one has almost surely*

$$\limsup_{n \rightarrow +\infty} \left(\frac{E_{\text{SLACK}}(n, m, \mathcal{E}_\lambda) - \text{OPT}}{\text{OPT}} \times \frac{n}{m^2 \ln \ln n} \right) \leq 1.$$

The proof is very similar to the one of Corollary 3.

6 Simulations

We first present the simulation setting in Section 6.1, before studying the δ_j 's and β_j 's in Section 6.2, and the energy consumption in Section 6.3.

6.1 Experimental Setting

All the following experiments have been conducted on Python 3.8.10. Two types of instances have been used. Both instances have in common that the platform is composed of $m = 100$ processors.

Theoretical instances have been generated using the *random* package. These instances have been generated following commonly used random distributions: the uniform distribution, $\mathcal{U}[0, 1]$; the exponential distribution of rate 1, $\mathcal{E}_1$; the distribution of cumulative distribution function $F(x) = x^\alpha$ where $0 < \alpha < \infty$ [12]. These simple distributions correspond to the ones for which there exist convergence results in the literature and they cover a wide range of situations.

Realistic instances have been generated using the experimental cumulative distribution functions of actual workloads [11]. These real workloads can be found on the Parallel Workload Archive from the website <https://www.cs.huji.ac.il/labs/parallel/workload/>. We used 3 specific instances: **KIT ForHLR II** with 114,355 tasks; **NASA Ames iPSC/860** with 18066 tasks; and **San Diego Supercomputer Center (SDSC) DataStar** with 84907 tasks.

6.2 Simulations: Study of δ_j and β_j

In this section, we describe the results of our simulations comparing the values of δ_j and β_j (the largest differences between the execution times of the processors and the tasks, as defined in Section 3) over the execution of SLACK and LPT.

In Figures 1 and 2, we can see the evolution of the quantities studied in Section 4 when bounding the performance of SLACK. The quantities are:

- β_j the difference between the largest and the shortest task of pack j (i.e., at step j of the algorithm), it describes the imbalance between consecutive tasks during the execution of the algorithms;

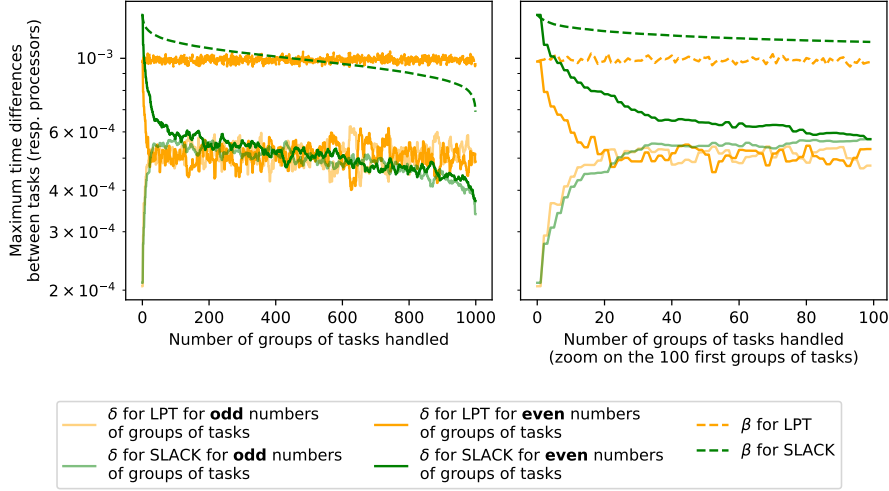


Figure 1: Evolution of δ_j and β_j (as defined in Section 3) during the execution of SLACK and LPT with the uniform distribution $\mathcal{U}[0, 1]$ for the tasks. Each execution is done with $m = 100$ processors and $n = 100\,000$ tasks. The right graph is a zoomed version of the 100 first values of δ_j and β_j . Each point represents the average value of δ_j (resp. β_j) over 30 executions.

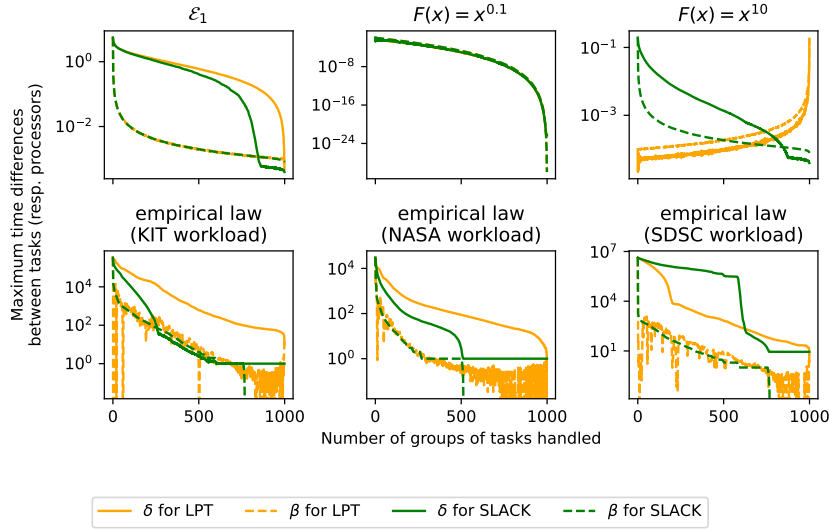


Figure 2: Evolution of δ_j and β_j (as defined in Section 3) during the execution of SLACK and LPT with various probability distributions for the tasks. Each execution is done with $m = 100$ processors and $n = 100\,000$ tasks. Each point represents the average value of δ_j (resp. β_j) over 30 executions.

- δ_j the difference between the largest processor and the shortest processor after step j of the algorithm (i.e., after allocating $j \times m$ tasks), it describes the imbalance between processors during the execution of the algorithms.

With these experiments, we can both investigate the relation we stated in Section 4, and investigate the unexplained “wave pattern” presented in [5].

In the case of tasks drawn through a uniform distribution with Figure 1, we observe that δ_j , the imbalance between processors, alternates between high and low values, in a sort of wave pattern. With a new representation of the pattern, we now present more elements explaining it. This pattern can be explained by the fact that the imbalance created by m consecutive tasks is then canceled by the m following tasks, as they have similar relative differences. Once the imbalance on the processors have decreased, the next m tasks will restore a new but smaller imbalance.

For most other distributions, on Figure 2, SLACK and LPT perform similarly in terms of makespan, which is characterized by the last value $\delta_{\frac{n}{m}}$. Out of our six examples, the only distribution for which SLACK performs significantly better than LPT is the distribution with cumulative distribution function $F(x) = x^{10}$, namely the one for which there are a few small tasks but many large ones.

A closer look at the evolution of δ_j and β_j gives more insights about the differences of execution between SLACK and LPT, and allows us to understand why SLACK performs better than LPT in some cases. Generally speaking, SLACK balances the different processors more quickly than LPT, and then keeps them balanced. In the specific case of $F(x) = x^{10}$, LPT performs significantly worse than SLACK because there is a high density of big tasks, and a low density of small tasks. It means that the big tasks are easy to balance whereas the small tasks are very different from each other. LPT finishes its execution with small tasks that have a very high difference β_j , whereas SLACK is able to balance the processors using big tasks.

6.3 Simulations: Energy minimization

In this section, we describe the results of the simulations, evaluating the energy consumed by the schedules of the algorithms derived from LPT and SLACK (as defined in Section 3).

We do not directly consider the energy E found by an algorithm because the value of W can vary a lot depending on the instance. Instead, we consider the relative difference between the energy found by the algorithm and a lower bound on OPT, i.e., $\frac{E - E_{\frac{W}{m}, \dots, \frac{W}{m}}}{E_{\frac{W}{m}, \dots, \frac{W}{m}}}$. We have shown in the proof of Theorem 2 that $E_{\frac{W}{m}, \dots, \frac{W}{m}}$ was indeed a lower bound on OPT.

The main conclusion that we can get from Figures 3 and 4 is that LPT and SLACK both perform very well on all created instances, both theoretical and realistic. The schedule that the two algorithms output is at most a few percents away from the optimal for very small instances, and the room for improvement rapidly decreases to less than $10^{-8}\%$ for larger instances.

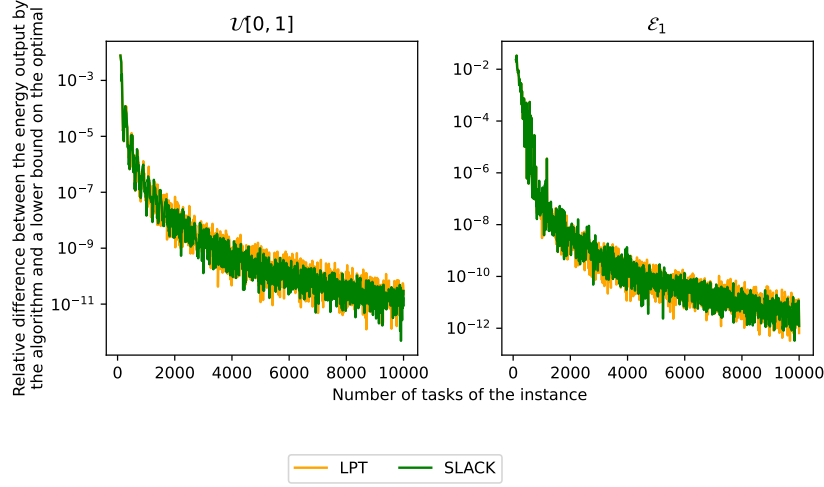


Figure 3: Relative difference between the energy found by SLACK or LPT with the speed strategy described in Theorem 2 and a lower bound on OPT, with various theoretical probability distributions for the tasks. Each execution is done with $m = 100$ processors. Each point represents the average value of energy over 30 executions.

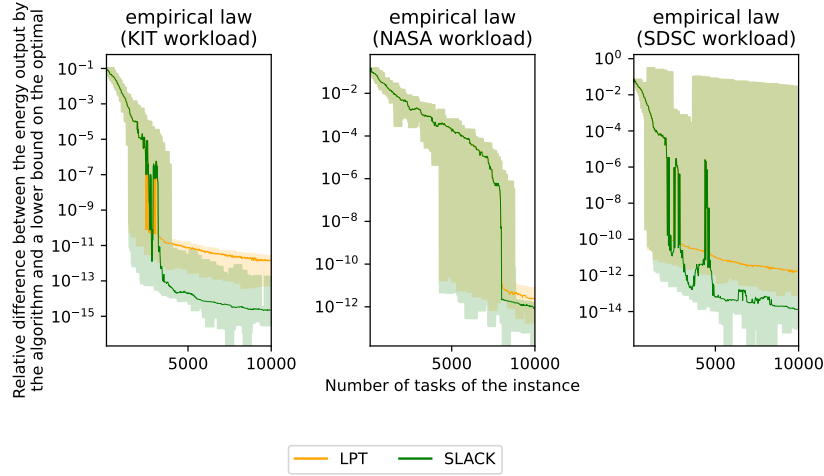


Figure 4: Relative difference between the energy found by SLACK or LPT with the speed strategy described in Theorem 2 and a lower bound on OPT, with various empirical probability distributions for the tasks. For each number of tasks, the execution is repeated 30 times with $m = 100$ processors. The thick lines represent the moving median, while the ribbons extend to the moving minimum and maximum over 45 values.

It can be noted that SLACK performs better than LPT on average, even if they are both near optimal.

7 Conclusion

The optimization of parallel computing platforms, especially for energy purposes, is a challenging societal issue. In this paper, we focus on SLACK, a recent heuristic that proves to be very efficient in practice for scheduling independent tasks on homogeneous machines. We have given, to the best of our knowledge, the first asymptotic performance results on the makespan for SLACK for tasks distributed randomly either uniformly or according to an exponential distribution. We have also shown how to adapt the numerous algorithms aiming at optimizing the makespan into algorithms dedicated to the optimization of the energy consumption. Based on SLACK, we were able to derive asymptotic energy performance results for both uniformly distributed tasks and exponentially distributed tasks. All these results exploit a common bound based of the shift between the most loaded and the least loaded processor during the execution of the heuristics. The experimental part of the paper proposes an empirical comparison of these shifts for SLACK and LPT. Finally, the performance for the energy problem is also studied experimentally, and both SLACK and LPT are shown to be near optimal in terms of energy consumption. In the future, it would be interesting to explore other energy cost models. It would also be interesting to study the bi-objective problem that considers both the makespan and the energy consumption. Finally, monitoring the energy consumption of schedules on actual machines using the proposed algorithms should be explored.

References

- [1] I. Bairamov, A. Berred, and A. Stepanov. Limit results for ordered uniform spacings. *Statistical Papers*, 51(1):227–240, 2010.
- [2] Mario Bambagini, Mauro Marinoni, Hakan Aydin, and Giorgio Buttazzo. Energy-aware scheduling for real-time systems: A survey. *ACM Trans. Embed. Comput. Syst.*, 15(1), January 2016.
- [3] Nikhil Bansal, Tracy Kimbrel, and Kirk Pruhs. Speed scaling to manage energy and temperature. *Journal of the ACM*, 54(1):1–39, 2007.
- [4] Luca Benini, Alessandro Bogliolo, and Giovanni De Micheli. A survey of design techniques for system-level dynamic power management. *IEEE Trans. on very large scale integration (VLSI) systems*, 8(3):299–316, 2000.
- [5] A. Benoit, L.-C. Canon, R. Elghazi, and P.-C. Héam. Update on the Asymptotic Optimality of LPT. In *Euro-Par 2021: Parallel Processing*, pages 55–69, 2021.

- [6] E. G. Coffman Jr, G. S. Lueker, and A. H. G. Rinnooy Kan. Asymptotic methods in the probabilistic analysis of sequencing and packing heuristics. *Management Science*, 34(3):266–290, 1988.
- [7] F. Della Croce and R. Scatamacchia. The longest processing time rule for identical parallel machines revisited. *Journal of Scheduling*, 23(2):163–176, 2020.
- [8] Pawel Czarnul, Jerzy Proficz, and Adam Krzywaniak. Energy-aware high-performance computing: survey of state-of-the-art tools, techniques, and environments. *Scientific Programming*, 2019, 2019.
- [9] L. Devroye. Laws of the Iterated Logarithm for Order Statistics of Uniform Spacings. *The Annals of Probability*, 9(5):860–867, 1981.
- [10] L. Devroye. The largest exponential spacing. *Utilitas Mathematica*, 25:303–313, 1984.
- [11] D.G. Feitelson, D. Tsafir, and D. Krakov. Experience with using the parallel workloads archive. *J. of Parallel and Distributed Comp.*, 74(10):2967–2982, 2014.
- [12] Johannes Bartholomeus Gerardus Frenk and A. H. G. Rinnooy Kan. The rate of convergence to optimality of the LPT rule. *Discrete Applied Mathematics*, 14(2):187–197, 1986.
- [13] R.L. Graham, E.L. Lawler, J.K. Lenstra, and AHG R. Kan. Optimization and approximation in deterministic sequencing and scheduling: a survey. *Annals of discrete mathematics*, 5:287–326, 1979.
- [14] Ronald L. Graham. Bounds on multiprocessing timing anomalies. *SIAM journal on Applied Mathematics*, 17(2):416–429, 1969.
- [15] X. Lin, Y. Wang, Q. Xie, and M. Pedram. Task scheduling with dynamic voltage and frequency scaling for energy minimization in the mobile cloud computing environment. *IEEE Trans. on Services Computing*, 8(2):175–186, 2014.
- [16] R. Loulou. Tight bounds and probabilistic analysis of two heuristics for parallel processor scheduling. *Mathematics of Operations Research*, 9(1):142–150, 1984.
- [17] Nanda Piersma and H Edwin Romeijn. Parallel machine scheduling: A probabilistic analysis. *Naval Research Logistics (NRL)*, 43(6):897–916, 1996.
- [18] Iosif Pinelis. Order statistics on the spacings between order statistics for the uniform distribution. *arXiv preprint arXiv:1909.06406*, 2019.

- [19] Ankit Thakkar, Kinjal Chaudhari, and Monika Shah. A comprehensive survey on energy-efficient power management techniques. *Procedia Computer Science*, 167:1189–1199, 2020.
- [20] Mark Weiser, Brent Welch, Alan Demers, and Scott Shenker. Scheduling for reduced CPU energy. In *Mobile Computing*, pages 449–471. Springer, 1994.
- [21] F. Yao, A. Demers, and S. Shenker. A scheduling model for reduced CPU energy. In *Proc. of IEEE 36th annual found. of computer science*, pages 374–382, 1995.



**RESEARCH CENTRE
GRENOBLE – RHÔNE-ALPES**

Inovallée
655 avenue de l'Europe Montbonnot
38334 Saint Ismier Cedex

Publisher
Inria
Domaine de Voluceau - Rocquencourt
BP 105 - 78153 Le Chesnay Cedex
inria.fr

ISSN 0249-6399