



Laboratoire de l'Informatique du Parallélisme

École Normale Supérieure de Lyon
Unité Mixte de Recherche CNRS-INRIA-ENS LYON n° 5668



***Optimal algorithms for scheduling divisible
workloads on heterogeneous systems***

Olivier Beaumont,
Arnaud Legrand,
Yves Robert

October 2002

Research Report N° 2002-36



École Normale Supérieure de Lyon

46 Allée d'Italie, 69364 Lyon Cedex 07, France

Téléphone : +33(0)4.72.72.80.37

Télécopieur : +33(0)4.72.72.80.80

Adresse électronique : lip@ens-lyon.fr



Optimal algorithms for scheduling divisible workloads on heterogeneous systems

Olivier Beaumont, Arnaud Legrand, Yves Robert

October 2002

Abstract

In this paper, we discuss several algorithms for scheduling divisible loads on heterogeneous systems. Our main contributions are (i) new optimality results for single-round algorithms and (ii) the design of an asymptotically optimal multi-round algorithm. This multi-round algorithm automatically performs resource selection, a difficult task that was previously left to the user. Because it is periodic, it is simpler to implement, and more robust to changes in the speeds of processors or communication links. On the theoretical side, to the best of our knowledge, this is the first published result assessing the absolute performance of a multi-round algorithm. On the practical side, extensive simulations reveal that our multi-round algorithm outperforms existing solutions on a large variety of platforms, especially when the communication-to-computation ratio is not very high (the difficult case).

Keywords: scheduling, divisible tasks, multi-round algorithms, asymptotical optimality

Résumé

Dans ce rapport, nous comparons un certain nombre d'algorithmes d'ordonnancement de tâches divisibles sur une plateforme hétérogène. Nos contributions principales sont (i) de nouveaux résultats d'optimalité pour les algorithmes à une étape et (ii) la conception d'un algorithme multi-étapes asymptotiquement optimal. Ce dernier algorithme effectue automatiquement la sélection des ressources à utiliser, tâche délicate généralement laissée à l'utilisateur. En raison de sa périodicité, il est plus facile à mettre en œuvre et plus robuste aux variations de charge des processeurs ou des liens de communications. D'un point de vue théorique, c'est, à notre connaissance, le premier résultat garanti sur les performances d'un algorithme multi-étapes. D'un point de vue plus appliqué, les simulations que nous avons menées montrent que cet algorithme est meilleur que les autres algorithmes sur une large variété de plateformes, tout particulièrement quand le rapport entre les communications et le calcul est élevé (le cas le plus délicat).

Mots-clés: ordonnancement, tâches divisibles, algorithmes multi-étapes, optimalité asymptotique

1 Introduction

Scheduling computational tasks on a given set of processors is a key issue for high-performance computing. In this paper, we restrict our attention to the processing of independent tasks whose size (and number) are a parameter of the scheduling algorithm. This corresponds to the divisible load model which has been widely studied in the last several years, and popularized by the landmark book written by Bharadwaj, Ghose, Mani and Robertazzi [6]. A divisible job is a job that can be arbitrarily split in a linear fashion among any number of processors. This corresponds to a perfectly parallel job: any sub-task can itself be processed in parallel, and on any number of processors. The applications of the divisible load model encompass a large spectrum of scientific problems, including among others Kalman filtering [22], image processing [19], video and multimedia broadcasting [1, 2], database searching [14, 7], and the processing of large distributed files [23] (see [6] for more examples).

On the practical side, the divisible load model provides a simple yet realistic framework to study the mapping on independent tasks on heterogeneous platforms. The granularity of the tasks can be arbitrarily chosen by the user, thereby providing a lot of flexibility in the implementation tradeoffs. On the theoretical side, the success of the divisible load model is mostly due to its analytical tractability. Optimal algorithms and closed-form formulas exist for the simplest instances of the divisible load problem. This is in sharp contrast with the theory of task graph scheduling, which abounds in NP completeness theorems [16, 15] and in inapproximability results [12, 3].

In this paper, the target computing platform is a heterogeneous master/worker platform, with p worker processes running on p processors labeled $P_1, P_2, \dots, P_p$. The master P_0 sends out chunks to workers over a network: we can think of a star-shaped network, with the master in the center. The master uses its network connection in exclusive mode: it can communicate with a single worker at any time-step. There are different scenarios for the workers, depending whether they can compute while receiving from the master (full overlap) or not. The overlap model is widely used in the literature, because it seems closer to the actual characteristics of state-of-the-art computing resources (but we point out that our results extend to both models, with and without overlap). For each communication of size L between the master and a worker, say P_i , we pay a latency g_i and a linear term $L \cdot G_i$, where G_i is the inverse of the bandwidth of the link between the master P_0 and P_i . In the original model of [6], all the latencies g_i are equal to zero, hence a linear cost model. However, latencies play an important role in current architectures [13], and more realistic models use the affine cost $g_i + LG_i$ for a message of size L . Finally, note that when $g_i = g$ and $G_i = G$ for $1 \leq i \leq p$, the star network can be viewed as a bus oriented network [22].

The master processor can distribute the chunks to the workers in a single round, (also called installment in [6]), so that there will be a single communication between the master and each worker. This is the simplest situation, but surprisingly the optimal solution for a heterogeneous star network is not known, even for a linear cost model. We provide the optimal solution in Section 4, thereby extending the results of [22] for bus oriented networks to heterogeneous platforms.

For large workloads, the single round approach is not efficient, because of the idle time incurred by the last processors to receive their chunks. To minimize the makespan, i.e. the total execution time, the master will send the chunks to the workers in multiple rounds: the communications will be shorter (less latency) and pipelined, and the workers will be able to compute the current chunk while receiving data for the next one. Deriving an efficient

solution becomes a challenging problem: how many rounds should be scheduled? what is the best size of the chunks for each round? Intuitively, the size of the chunks should be small in the first rounds, so as to start all the workers as soon as possible. Then the chunk size should increase to a steady-state value, to be determined so as to optimize the usage of the total available bandwidth of the network. Finally the chunk size should be decreased while reaching the end of the computation. In Chapter 10 of [6], there is no quantified value provided for the number of rounds to be used. Recently, Altilar and Paker [1, 2], and Yang and Casanova [24] have introduced multi-round algorithms and analytically expressed their performance. We discuss these algorithms, and others, in Section 3, which is devoted to related work. To the best of our knowledge, no optimality result has ever been obtained for multi-round algorithms on heterogeneous platforms. The most important result of this paper is to fill this gap: in Section 5, we design a periodic multi-round algorithm and we establish its asymptotic optimality.

The rest of the paper is organized as follows. We begin with models for computation and communication costs in Section 2. Next we review related results in Section 3. Then we deal with single-round algorithms in Section 4. We proceed to multi-round algorithms in Section 5. We provide some simulations in Section 6. Finally, we state some concluding remarks in Section 7.

2 Models

As already said, we assume a total workload W_{total} that is perfectly divisible into an arbitrary number of pieces, or chunks. Usually, it is assumed that the master itself has no processing capability, because otherwise we can add a fictitious extra worker paying no communication cost to simulate the master. There is a wide acceptance in the literature on using linear costs to model computation costs. Worker P_i will require $\alpha_i w_i$ time-units to process a chunk of size α_i . However, Yang and Casanova [24] suggest to add a start-up cost, or computation latency, so that the cost becomes $z_i + \alpha_i w_i$ for P_i to process a chunk of size α_i ; they emphasize the importance of adding such a latency to obtain realistic results in some data-sweep applications [9]. In the following, we stick to linear computational costs, but we will later mention which equations to modify to take latencies into account in the multi-round algorithms.

Modeling communication costs is more difficult, and several models have been proposed. In the original approach [6], communication costs were also assumed linear. The master would need $\alpha_i G_i$ time-units to send a chunk of size α_i to P_i . While acceptable for large messages, the model becomes quite unrealistic for small messages. For instance in [6], the authors recognize that infinitely small messages would be the best solution for multi-round algorithms with this crude model. Communication latencies g_i have been introduced by Drozdowski [14] and are now widely used¹: the master needs $g_i + \alpha_i G_i$ to send a chunk of size α_i to worker P_i . An even more accurate model has been proposed by Rosenberg [21] and further investigated by Yang and Casanova [24]. They suggest to use the expression $g'_i + \alpha_i G_i + g''_i$, where the first latency g'_i is not overlappable, while the second latency g''_i is overlappable with the next communication. The master may send another message $g'_i + \alpha_i G_i$ time-units later, while the worker cannot start computing before $g'_i + \alpha_i G_i + g''_i$ time-units. The overlappable latency was introduced to model pipelined networking. Again, we restrict ourselves to non-overlappable latencies, but we will indicate how to incorporate them in the design of multi-round algorithms. Finally,

¹Because there is no consensus on the notations, we borrowed the notations g_i and G_i from Wang et al [23]

note that other models [18, 17] assume a fixed communication cost to dispatch chunks of any size, which seems much less realistic than adopting an affine expression with a startup and a linear term proportional to the chunk size.

Next, there is to discuss the amount of computation and communication that can be overlapped. In the model *with overlap*, each worker is capable of receiving the next chunk from the master while computing the current chunk. This corresponds to workers *equipped with a front end* in [6]. In the *no overlap* model, each worker executes communications and computations sequentially. Of course this distinction of models only applies to multi-round algorithms, because in a single-round algorithm it is impossible to overlap the communication with independent computation. When dealing with multi-round algorithms in Section 5, we will elaborate results for both models, with and without overlap..

The last question is the number of simultaneous communications that can be handled by the master. With few exceptions, the *one-port model* is assumed: the master can communicate with at most a worker at a given time-step (except may be for the short time-slice corresponding to the overlappable latency). However, as pointed out by Yang and Casanova [24], the one-port model is nicely suited to LAN network connections but a multi-port model could be used for WAN network connections.

In conclusion, we retain the following model:

1. One-port for the master (at most one communication to a worker at any time-step)
2. Communication-computation overlap for the workers
3. Linear computation costs $\alpha_i w_i$ for a chunk of size α_i processed by P_i , $1 \leq i \leq p$
4. Affine communication cost $g_i + \alpha_i G_i$ to send a chunk of size α_i from P_0 to worker P_i , $1 \leq i \leq p$

We discuss some extensions of this model when dealing with multi-round algorithms.

3 Related results

We divide this overview into two categories: results for single-round algorithms, and results for multi-round algorithms. We restrict ourselves to master/worker platforms, which includes bus-oriented and star-shaped networks. See [6] for results on processor trees and [14] for hypercubes.

3.1 Single-round algorithms

For single-round algorithms, the first problem is to determine in which order the chunks should be sent to the different workers. Since the master can handle only one communication at a given time step, the solution is as depicted in Figure 1. Once the communication order has been determined, the second problem is to decide how much work should be allocated to each processor P_i . The final objective is to minimize the makespan, i.e. the total execution time.

In the case of a homogeneous (bus-oriented) platform (all G_i are equal to G), and using a linear cost model for computation (all g_i are equal to zero), Bataineh, Hsiung and Robertazzi [4, 22] have derived an optimal solution, together with closed-form expressions for the

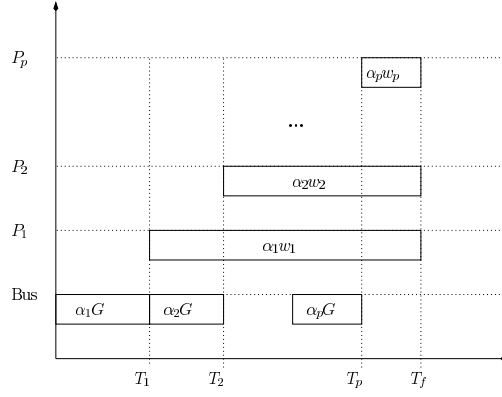


Figure 1: Pattern of a solution for dispatching the load of a divisible job, using a bus-oriented platform ($G_i = G$) All workers complete execution at the same time-step T_f .

makespan T_f . This solution is surprisingly simple. Let α_i denote the fraction of workload assigned to worker P_i , where $\sum_{i=1}^p \alpha_i = W_{\text{total}}$, and let T_i denote the time elapsed before P_i begins its processing. Thus, $T_f = \max_i (T_i + \alpha_i w_i)$.

First, one can prove that all the processors must finish their work at the same time (i.e. $T_i + \alpha_i w_i = T_f$, $\forall i$). Indeed, otherwise, some work could be transferred from a busy processor to an idle one in order to reduce T_f . Thus, the following system of equation holds,

$$\begin{cases} T_f - T_i = \alpha_i w_i, & \forall 1 \leq i \leq p \\ T_{i+1} - T_i = \alpha_{i+1} G & \forall 1 \leq i \leq p-1 \end{cases}$$

if data is sent successively to $P_1, \dots, P_p$. Closed forms can be obtained for both the α_i 's and T_f . These closed form are rather complicated, although the method for obtaining them is elementary, and we refer the reader to [22] to find the actual algebraic expressions. The surprising and interesting point is that the overall computational time T_f does not depend upon the order chosen for sending data to the different processors, so that the ordering $P_1, \dots, P_p$ is in fact optimal.

Later, Charcranoon, Roberatzzi and Luryi [11] have partially extended this work to heterogeneous (star-shaped) platforms: they still use linear communication costs, but with different G_i 's. The results are less satisfying than in the case of the bus. Indeed, the main known result is that if data is sent to the different processors in a given order (say, again, $P_1, \dots, P_p$), then closed forms can be obtained for both the α_i 's and T_f . Unfortunately, the makespan T_f strongly depends on the communication ordering, and the result stating that all the processors must finish their work at the same time-step is no longer valid for all communication orderings. To the best of our knowledge, the optimal communication ordering is not known, and we provide the optimal solution in Section 4.

Moving to affine communication costs rather than linear communication costs, several results have been published, among others [19, 14, 7, 21]. In 1997, Drozdowski [14] stated that the complexity of determining the optimal makespan for a general star-shaped platform (different g_i 's and different G_i 's) is not known, and to the best of our knowledge the problem is still open. We point out that Drozdowski [14] proposes an interesting mixed linear programming formulation of the problem. In the following program, $x_{i,j}$ is a boolean variable that equals 1 if P_j is chosen for the i -th communication from the master:

$$\begin{array}{ll}
\text{MINIMIZE } T_f, \\
\text{SUBJECT TO} \\
\left\{ \begin{array}{ll}
(1) \alpha_i \geq 0 & 1 \leq i \leq p \\
(2) \sum_{i=1}^p \alpha_i = W_{\text{total}} & \\
(3) x_{i,j} \in \{0, 1\} & 1 \leq i, j \leq p \\
(4) \sum_{i=1}^p x_{i,j} = 1 & 1 \leq j \leq p \\
(5) \sum_{j=1}^p x_{i,j} = 1 & 1 \leq i \leq p \\
(6) \sum_{i=1}^p x_{1,i}(g_i + \alpha_i G_i + \alpha_i w_i) \leq T_f & \text{(first communication)} \\
(7) \sum_{k=1}^{j-1} \sum_{i=1}^p x_{k,i}(g_i + \alpha_i G_i) & 2 \leq j \leq p \text{ (} j\text{-th communication)} \\
\quad + x_{j,i}(g_i + \alpha_i G_i + \alpha_i w_i) \leq T_f &
\end{array} \right.
\end{array}$$

Equation (4) states that exactly one processor is activated for the j -th communication, and equation (5) states that each processor is activated exactly once. Equation (6) is a particular case of equation (7), which expresses that the processor selected for the j -th communication (where $j = 1$ in equation (6) and $j \geq 2$ in equation (7)) must wait for the previous communications to complete before its own communication and computation, and that all this quantity is a lower bound of the makespan. As pointed out by Drozdowski [14], this mixed linear program may have no solution if all the workers are not involved in the optimal solution (it may well be the case that using a strict subset of the resources proves more efficient), so the formulation is not fully general.

3.2 Multi-round algorithms

Several multi-round algorithms have been proposed in the literature [6, 1, 2, 24] but in general they have been validated through simulations or experiments rather than with analytical formulas. This is not surprising: deriving the adequate number of rounds is a challenging task. On one hand short rounds minimize idle times in the beginning, and enable to better overlap computations and communications. On the other hand longer rounds mean less latency overheads.

Technically, a round is defined as a sequence of communications to different workers, one per worker, and deciding whether to use all workers or a strict subset of the workers is a difficult question. Even worse, should a strict subset be used, there is no reason for the subset to remain the same from one round to another.

Let $W^{(k)}$ be the total size of the chunks assigned to the workers during round k : $W^{(k)} = \sum_{i=1}^p \alpha_i^{(k)}$, where $\alpha_i^{(k)}$ is the chunk size of P_i at round k . Intuitively, $W^{(k)}$ should be small for the first rounds, then reach an adequate value, and then decrease in the last rounds. Yang and Casanova [24] propose that $W^{(k)}$ follows a geometric progression, and within each round that all involved processors compute for the same amount of time². These simplifying assumptions enable them to derive analytical expressions for the total execution time, and the optimal number of rounds is then derived through some numerical optimization technique. The results are technically involved but very interesting. However, there remains two main limitations to this approach: (i) resource selection (determining the best subset) is performed heuristically, and (ii) there is no fundamental reason to privilege a geometric progression for the round sizes, any other monotonic and sufficiently “regular” function could be adopted.

²The geometric progression is stopped when approaching the end of the execution, so that all processors terminate working simultaneously.

In Section 5, we derive a periodic algorithm which is asymptotically optimal. This algorithm is simple, because rounds are repeated identically one after the others. The key-issues, i.e. the optimal number of chunks, resource selection and chunk size assignments within a chunk, are all solved through a relaxed linear program in rational numbers (hence a low-degree polynomial complexity).

4 New results for single-round algorithms

In this section, we propose a new proof method for the optimal distribution of the work to the processors in single-round algorithms. This approach enables us to retrieve some well known results, and to establish new ones.

The approach is based upon the comparison of the amount of work that is performed by the first two workers. To simplify notations, assume that P_1 and P_2 have been selected as the first two workers. There are two possible orderings, as illustrated in Figure 2. For each ordering, we will determine the total number of tasks $\alpha_1 + \alpha_2$ that have been processed in T time-units, and the total occupation t_2 of the communication medium during this time interval. We denote with upper-script (A) (resp. (B)) all the quantities related to the first (resp. second) ordering.

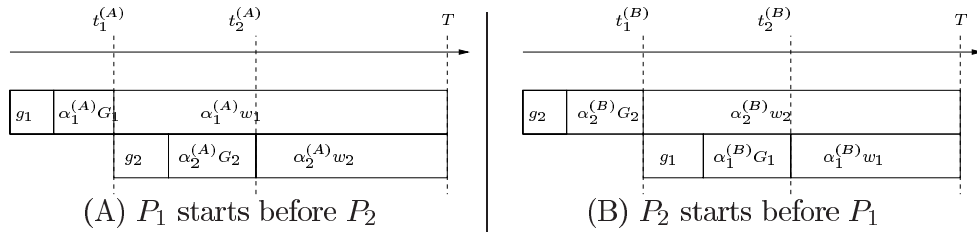


Figure 2: Comparison of the two possible orderings.

Let us first determine the different quantities $\alpha_1^{(A)}$, $\alpha_2^{(A)}$, $t_1^{(A)}$ and $t_2^{(A)}$ for the left ordering in Figure 2:

- $g_1 + \alpha_1^{(A)}(G_1 + w_1) = T \implies T = \frac{T - G_1}{G_1 + w_1}$.
- $t_1^{(A)} = g_1 + \alpha_1^{(A)}G_1 \implies t_1^{(A)} = \frac{g_1 w_1 + T G_1}{G_1 + w_1}$.
- $t_1^{(A)} + g_2 + \alpha_2^{(A)}(G_2 + w_2) = T \implies \alpha_2^{(A)} = \frac{T w_1 - g_2(G_1 + w_1) - g_1 w_1}{(G_1 + w_1)(G_2 + w_2)}$.
- $t_2^{(A)} = t_1^{(A)} + g_2 + \alpha_2^{(A)}G_2 \implies t_2^{(A)} = \frac{T(G_2 G_1 + G_2 w_1 + G_1 w_2) + (g_1 + g_2)w_1 w_2 + G_1 g_2 w_1}{(G_1 + w_1)(G_2 + w_2)}$.

Therefore, the overall number of processed tasks is equal to

$$\alpha_1^{(A)} + \alpha_2^{(A)} = \frac{T - g_1}{G_1 + w_1} + \frac{T - g_2}{G_2 + w_2} - \frac{T G_1 + g_1 w_1}{(G_1 + w_1)(G_2 + w_2)},$$

and the overall occupation time of the network medium is equal to

$$t_2^{(A)} = \frac{T(G_2 G_1 + G_2 w_1 + G_1 w_2) + (g_1 + g_2)w_1 w_2 + G_1 g_2 w_1}{(G_1 + w_1)(G_2 + w_2)}.$$

These expressions are rather complicated. Nevertheless, it is possible to obtain simple expressions when expressing the differences between situation (A) and situation (B). Indeed, we have:

$$\left(\alpha_1^{(A)} + \alpha_2^{(A)}\right) - \left(\alpha_1^{(B)} + \alpha_2^{(B)}\right) = \frac{g_2 w_2 - g_1 w_1 + T(G_2 - G_1)}{(G_1 + w_1)(G_2 + w_2)}$$

and

$$t_2^{(A)} - t_2^{(B)} = \frac{G_1 g_2 w_2 - G_2 g_1 w_1}{(G_1 + w_1)(G_2 + w_2)}.$$

Thanks to these expressions, we can derive the optimal distribution in some special cases:

1. $g_1 = g_2 = 0$: Then, the occupation of the communication medium does not depend on the communication ordering, since $t_2^{(B)} = t_2^{(A)}$. Therefore, we only need to consider the number of processed tasks in both situations. Since

$$\left(\alpha_1^{(A)} + \alpha_2^{(A)}\right) \geq \left(\alpha_1^{(B)} + \alpha_2^{(B)}\right) \iff G_2 \geq G_1,$$

we have better to send tasks to the processor with the smallest G_i first.

In the case of p processors, we sort them so that $G_1 \leq G_2 \leq \dots \leq G_p$. We state the first result:

Theorem 1. *When all communication latencies g_i are equal to 0, sort the p processors so that $G_1 \leq G_2 \leq \dots \leq G_p$. Then the ordering where tasks are sent to $P_1, P_2, \dots, P_p$ is optimal.*

Proof. Consider an optimal ordering of the communications σ , where tasks are sent successively to $P_{\sigma(1)}, P_{\sigma(2)}, \dots, P_{\sigma(p)}$. Let us denote by i , if it exists, the smallest index satisfying $\sigma(i) > \sigma(i+1)$. Let us consider the following ordering:

$$P_{\sigma(1)}, \dots, P_{\sigma(i-1)}, P_{\sigma(i+1)}, P_{\sigma(i)}, P_{\sigma(i+2)}, \dots, P_{\sigma(p)}.$$

Then, $P_{\sigma(1)}, \dots, P_{\sigma(i-1)}, P_{\sigma(i+2)}, \dots, P_{\sigma(p)}$ perform exactly the same number of tasks, since the exchange does not affect the overall communication time, but together, $P_{\sigma(i+1)}$ and $P_{\sigma(i)}$ perform $\frac{T(G_2 - G_1)}{(G_1 + w_1)(G_2 + w_2)}$ more tasks, where T denotes the remaining time after communications to $P_{\sigma(1)}, \dots, P_{\sigma(i-1)}$. Since $G_{\sigma(i+1)} > G_{\sigma(i)}$, there exists an optimal ordering where tasks are sent accordingly to increasing values of the G_i 's. $\square$

Once the optimal ordering is known, we can use the formulas in [11] to derive the optimal assignment of works to processors, thereby filling the gap towards obtaining an optimal solution in the heterogeneous case. If all the G_i 's are equal, then we find the classical result of [4], stating that the number of processed tasks does not depend of the communication ordering.

2. $G_1 = G_2 = G$, but the start-up times g_1 and g_2 are different. Then,

$$\left(\alpha_1^{(A)} + \alpha_2^{(A)}\right) - \left(\alpha_1^{(B)} + \alpha_2^{(B)}\right) = \frac{g_2 w_2 - g_1 w_1}{(G + w_1)(G + w_2)},$$

and

$$t_2^{(A)} - t_2^{(B)} = G \frac{g_2 w_2 - g_1 w_1}{(G + w_1)(G + w_2)}.$$

Therefore, the number of tasks processed in situation (A) is larger if and only if $g_1w_1 \leq g_2w_2$, but in this case, the occupation of the communication medium is also larger. Nevertheless, since

$$\frac{(\alpha_1^{(A)} + \alpha_2^{(A)}) - (\alpha_1^{(B)} + \alpha_2^{(B)})}{t_2^{(B)} - t_2^{(A)}} = \frac{1}{G},$$

the amount of communication required per task is optimal for those extra tasks. Thus, in order to maximize the number of processed tasks, we have better to send tasks to the processor with the smallest g_iw_i first.

With p processors, we extend this result as follows:

Theorem 2. *When all elemental transfer times G_i are equal to G , sort the p processors so that $g_1w_1 \leq g_2w_2 \leq \dots \leq g_pw_p$. Then the ordering where tasks are sent to $P_1, P_2, \dots, P_p$ is optimal.*

Proof. Consider an optimal ordering of the communications σ , where tasks are sent successively to $P_{\sigma(1)}, P_{\sigma(2)}, \dots, P_{\sigma(p)}$. Denote by i , if it exists, the smallest index satisfying $\sigma(i) > \sigma(i+1)$. We can derive the following ordering:

$$P_{\sigma(1)}, \dots, P_{\sigma(i-1)}, P_{\sigma(i+1)}, P_{\sigma(i)}, P_{\sigma(i+2)}, \dots, P_{\sigma(p)}.$$

This is a solution that performs at least as many tasks than the optimal ordering. Since we exchanged $P_{\sigma(i)}$ and $P_{\sigma(i+1)}$, we are able to process β more tasks with $P_{\sigma(i)}$ and $P_{\sigma(i+1)}$, where $\beta = \frac{g_{\sigma(i)}w_{\sigma(i)} - g_{\sigma(i+1)}w_{\sigma(i+1)}}{(G+w_{\sigma(i)})(G+w_{\sigma(i+1)})}$, but that would induce an extra communication cost of βG , as depicted in Figure 3. Nevertheless, we can send (using the notations of Figure 3) $\alpha_{i+1}^{(C)}$ tasks to $P_{\sigma(i+1)}$ and $\alpha_i^{(C)}$ tasks to $P_{\sigma(i)}$, where

$$\alpha_{i+1}^{(C)} = \alpha_{i+1}^{(B)} - \frac{\beta}{2} \text{ and } \alpha_i^{(C)} = \alpha_i^{(B)} - \frac{\beta}{2}.$$

In this case, since the number of received tasks is lower for $P_{\sigma(i)}$ and $P_{\sigma(i+1)}$, both processors are able to complete their processing within time bounds. Moreover, the communication medium is occupied exactly the same time, and the number of processed tasks is the same as in the optimal solution. Thus, there exists an optimal ordering where tasks are sent accordingly to increasing values of the g_iw_i 's. $\square$

5 Asymptotically optimal multi-round algorithms

In this section, we derive asymptotically optimal algorithms for the multi-round distribution of divisible tasks, when slave processors are either able or not to overlap their processing with incoming communications.

5.1 No overlap

The sketch of the algorithm that we propose is as follows: the overall processing time T is divided into k regular periods of duration T_p (hence $T = kT_p$, but k (and T_p) are yet to be determined).

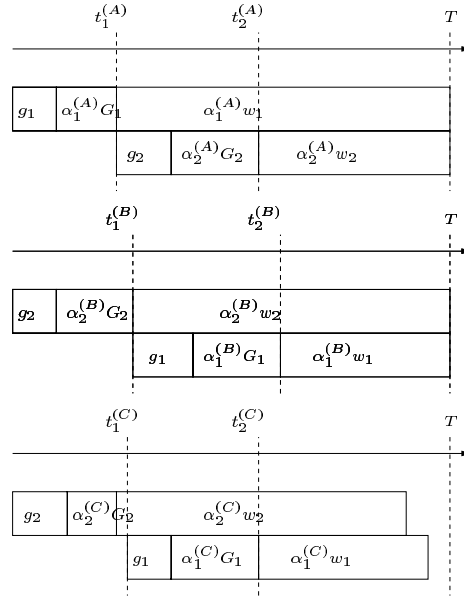


Figure 3: Illustrating the proof of Theorem 2.

During a period of duration T_p , the master processor sends α_i tasks to slave processor P_i (see Figure 4 for an example). It may well be the case that not all the processors are involved in the computation. Let $\mathcal{I} \subset \{1, \dots, p\}$ represent the subset of indices of participating processors. For all $i \in \mathcal{I}$, the α_i 's must satisfy the following inequality, stating that communication resources are not exceeded:

$$\sum_{i \in \mathcal{I}} (g_i + \alpha_i G_i) \leq T_p. \quad (1)$$

Since the processors cannot overlap communications and processing, the following inequalities also hold true:

$$\forall 1 \leq i \leq p, i \in \mathcal{I}, \quad g_i + \alpha_i (G_i + w_i) \leq T_p.$$

Let us denote by β_i the averaged number of tasks that slave P_i processes during one time unit, then the system becomes

$$\begin{cases} \forall 1 \leq i \leq p, i \in \mathcal{I}, & \beta_i (G_i + w_i) \leq 1 - \frac{g_i}{T_p} \\ \sum_{i \in \mathcal{I}} \beta_i G_i \leq 1 - \frac{\sum_{i \in \mathcal{I}} g_i}{T_p} \end{cases},$$

and our aim is to maximize the overall number of tasks processed during one time unit, i.e. $n_{\max} = \sum_{i \in \mathcal{I}} \beta_i$.

Let us consider the following linear program:

$$\begin{aligned} & \text{MAXIMIZE } \sum_{i=1}^p x_i, \\ & \begin{cases} \forall 1 \leq i \leq p, & x_i (G_i + w_i) \leq 1 - \frac{\sum_{i=1}^p g_i}{T_p} \\ \sum_{i=1}^p x_i G_i \leq 1 - \frac{\sum_{i=1}^p g_i}{T_p} \end{cases}. \end{aligned}$$

This linear program is more constrained than previous one, since $1 - \frac{g_i}{T_p}$ has been replaced by $1 - \frac{\sum_{i=1}^p g_i}{T_p}$ in p inequalities. Therefore, the objective value for $\sum_{i=1}^p x_i$ satisfies $\sum_{i=1}^p x_i \leq$

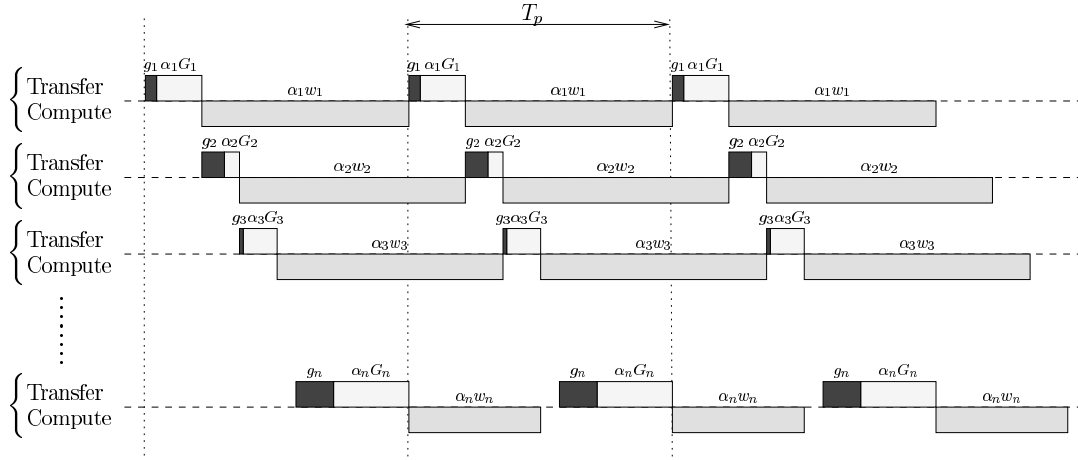


Figure 4: Sketch of a periodic multi-round schedule using the first n workers P_1 to P_n , where $n \leq p$.

$n_{\max}$. The linear program can be solved using a package similar to Maple [10] (we have rational numbers), but it turns out that the technique developed in [5] enables us to obtain the solution in closed form. We refer the reader to [5] for the complete proof. Let us sort the G_i 's so that $G_1 \leq G_2 \leq \dots \leq G_p$, and let q be the largest index so that $\sum_{i=1}^q \frac{G_i}{G_i + w_i} \leq 1$. If $q < p$, let ϵ denote the quantity $1 - \sum_{i=1}^q \frac{G_i}{G_i + w_i}$. The optimal solution to the linear program is obtained with

$$\forall 1 \leq i \leq q, \quad x_i = \frac{1 - \frac{\sum_{i=1}^p g_i}{T_p}}{G_i + w_i}$$

and (if $q < p$):

$$x_{q+1} = \left(1 - \frac{\sum_{i=1}^p g_i}{T_p}\right) \left(\frac{\epsilon}{G_{q+1}}\right),$$

and $x_{q+2} = x_{q+3} = \dots = x_p = 0$.

With these values, we obtain:

$$\sum_{i=1}^p x_i = \left(1 - \frac{\sum_{i=1}^p g_i}{T_p}\right) \left(\sum_{i=1}^p \frac{1}{G_i + w_i} + \frac{\epsilon}{G_{p+1}}\right).$$

Thus, we verify that $n_{\max}$ satisfies

$$n_{\max} \geq \left(1 - \frac{\sum_{i=1}^p g_i}{T_p}\right) \left(\sum_{i=1}^p \frac{1}{G_i + w_i} + \frac{\epsilon}{G_{p+1}}\right).$$

Let us denote by $n_{\max}^*$ the optimal number of tasks that can be processed within one unit of time. If we denote by β_i^* the optimal number of tasks that can be processed by slave P_i within one unit of time, the β_i^* 's satisfy the following set of inequalities, in which the g_i 's have been withdrawn:

$$\begin{cases} \forall 1 \leq i \leq p, & \beta_i^*(G_i + w_i) \leq 1 \\ \sum_{i=1}^p \beta_i^* G_i \leq 1 \end{cases}$$

Here, because we have no latencies, we can safely assume that all the processors are involved (and let $\beta_i^* = 0$ for some of them). We derive that:

$$n_{\max}^* \leq \left(1 - \frac{\sum_i g_i}{T_p}\right) \left(\sum_i^p \frac{1}{G_i + w_i} + \frac{\epsilon}{G_{q+1}}\right).$$

If we consider a large number B of tasks to be processed and if we denote by $T_{\max}^*$ the optimal time necessary to process them, then

$$T_{\max}^* \geq \frac{B}{n_{\max}^*} \geq \frac{B}{\left(\sum_{i=1}^p \frac{1}{G_i + w_i} + \frac{\epsilon}{G_{q+1}}\right)}.$$

Let us denote by T the time necessary to process all B tasks with the algorithm that we propose. Since the first period is lost for processing, then the number k of necessary periods satisfies $n_{\max} T_p (k - 1) \geq B$ so that we choose

$$k = \frac{B}{n_{\max} T_p} + 2.$$

Therefore,

$$T \leq \frac{B}{n_{\max}} + 2T_p \leq \frac{B}{\left(\sum_{i=1}^q \frac{1}{G_i + w_i} + \frac{\epsilon}{G_{q+1}}\right)} \left(\frac{1}{1 - \sum_{i=1}^p \frac{g_i}{T_p}}\right) + 2T_p,$$

and therefore, if $T_p \geq 2 \sum_{i=1}^p g_i$,

$$T \leq T_{\max}^* + 2 \sum_{i=1}^p g_i \frac{T_{\max}^*}{T_p} + 2T_p.$$

Finally, if we set $T_p = \sqrt{T_{\max}^*}$, we check that

$$T \leq T_{\max}^* + 2 \left(\sum_{i=1}^p g_i + 1\right) \sqrt{T_{\max}^*} = T_{\max}^* + O(\sqrt{T_{\max}^*}),$$

which achieves of proof of the asymptotic optimality of our algorithm.

Note that resource selection is part of our explicit solution to the linear program: to give an intuitive explanation of the analytical solution, processors are greedily selected, fast-communicating processors first, as long as the communication to communication-added-to-computation ratio is not exceeded.

Also, note that it is easy to include a computation latency z_i , as suggested by Yang and Casanova [24]: simply replace G_i by $G_i + z_i$ in the formulas.

We formally state our main result:

Theorem 3. *For arbitrary values of g_i , G_i and w_i , and assuming no communication-computation overlap, the previous periodic multi-round algorithm is asymptotically optimal. Closed-form expressions for resource selection and task assignment are provided by the algorithm, whose complexity does not depend upon the total amount of work to execute.*

5.2 With overlap

In the case where slaves are able to overlap communications and processing, the algorithm that we propose is very similar to the previous one. Thus, we do not detail the proof. During time period $i + 1$, the slave processors process the tasks that they have received during time period i , so that no processing occurs during the first period, and no communication occurs during the last period. The system of inequalities for one time unit using our algorithm becomes:

$$\begin{cases} \forall 1 \leq i \leq p, & i \in \mathcal{I} \quad \beta_i w_i \leq 1 \\ \sum_{i \in \mathcal{I}} \beta_i G_i \leq 1 - \frac{\sum_{i=1}^p g_i}{T_p} \end{cases}$$

and we can prove, as previously that $n_{\max}$ satisfies

$$n_{\max} \geq \left(1 - \frac{\sum_{i=1}^p g_i}{T_p}\right) \left(\sum_{i=1}^q \frac{1}{w_i} + \frac{\epsilon}{G_{q+1}}\right).$$

where q is the largest index so that $\sum_{i=1}^q \frac{G_i}{w_i} \leq 1$ and if $q < p$, $\epsilon = 1 - \sum_{i=1}^q \frac{G_i}{w_i}$. Similarly,

$$n_{\max}^* \leq \left(\sum_{i=1}^q \frac{1}{w_i} + \frac{\epsilon}{G_{q+1}}\right).$$

and we obtain

$$T \leq T_{\max}^* + 2\left(\sum_{i=1}^p g_i + 1\right)\sqrt{T_{\max}^*},$$

what achieves the proof of the asymptotic optimality of our algorithm.

To summarize, we still sort the processors according to the bandwidths G_i and we select them greedily until the sum of the ratios $\frac{G_i}{w_i}$ (instead of $\frac{G_i}{G_i + w_i}$) exceeds 1. We state the result formally:

Theorem 4. *For arbitrary values of g_i , G_i and w_i , and assuming full communication-computation overlap, the previous periodic multi-round algorithm is asymptotically optimal. Closed-form expressions for resource selection and task assignment are provided by the algorithm, whose complexity does not depend upon the total amount of work to execute.*

6 Simulations

In order to evaluate our multi-round algorithm, we have crafted a simulation with the SimGrid simulator [8, 20]. One major interest of relying on SimGrid is that all machine and network characteristics used in the simulations correspond to realistic values taken from the SimGrid database. We detail below the platforms that we have simulated.

In the experiments, we let the total workload size W_{total} vary in terms of workload units (or tasks) whose number range from 100 to 2000 by step of 100. Of course, the divisible load model applies here, so we assign fractional numbers of units to the processors. We let the size of a workload unit itself (i.e. the number of floating-point operations performed per unit) vary from one set of experiments to the other, so as to investigate different communication-to-computation ratios for a given application/platform pair.

In the experiments, no overlap of communications by computations was possible. We have compared our no-overlap multi-round algorithm with the multi-installment algorithm proposed in [6]. We have used a total of eleven heuristics. Three heuristics are different variants of the linear programming formulation, and the multi-installment algorithm has been tested for 1 to 8 installments. Here is a description of the eleven heuristics:

L.P with fixed period We use here the simplest variant of linear programming. We arbitrarily fix the value of the period to $T_p = 2000$. While there remains tasks to process, we allocate them to the workers according to a variant of Equation 1: we maximize $\sum_{i=1}^p \alpha_i$ subject to $\sum_{i=1}^p (g_i + \alpha_i G_i) \leq T_p$ and $g_i + \alpha_i(G_i + w_i) \leq T_p$ for all i , $1 \leq i \leq p$. The problem is slightly over-constrained, in that we include all p latencies in Equation 1 governing the bandwidth utilization, rather than only those of the participating processors $i \in \mathcal{I}$, as in the original formulation. Of course if the linear program returns $\alpha_i = 0$ for some i , we do not schedule the empty communication. This approximation is very good for large values of T_p , and provides a simple yet efficient task allocation if the period T_p is known a priori.

L. P. with fixed square-root period At the beginning of the computation, an evaluation of the optimal time T needed to process the whole set of task is computed, by neglecting all latencies. This works as follows: we assume a perfectly load-balance of the work and write $(G_i + w_i)\alpha_i = \text{Constant}$, with $\sum_{i=1}^p \alpha_i = W_{\text{total}}$. Once we have T , we let $T_p = \sqrt{T}$, and $\alpha_i = x_i T_p$, $1 \leq i \leq n$, where the x_i 's are the values computed in Section 5.1.

L. P. with adaptive period This is a slight modification of the previous heuristic. At each round, the period T_p is recomputed as $T_p = \sqrt{T}$, where T now is an estimation of the total time needed to process the remaining work units (rather than the total time for all units). In the very last steps of the heuristic, we stop the process and do not decrease T_p below the time needed to process the last work unit.

M.I.x This is the Multi-Installment procedure of [6] with x rounds. A set of linear equations, whose number of variables depends on x , is proposed in [6]. From these equations, it is possible to derive the amount of work units to distribute to each process at each round. Note that these equations do not take in account the latency.

6.1 Homogeneous platforms, no latency

The first set of experiments deals with homogeneous platforms, made up with PIII 1GHz processors (delivering 114.444 Mflops), interconnected through an Ethernet 100 Mbits/sec (but measured at 32.10 Mbits/sec bandwidth). The number of processors ranges from 1 to 20. One workload unit amounts to 1 GFlops of computations and 2 Mbits of data exchange. In other words, the workers communicate during 0.06 seconds to be able to compute during 9 seconds. The communication-to-computation ratio is quite low, which makes it easier to obtain good performances.

Figure 5 depicts the behavior of the eleven heuristics for a 5-processor platform. The general behavior is the same for other platform sizes. We see that L.P with fixed period and M.I.1, the one-round strategy, perform very badly. All the other strategies look similar. To outline the differences, we plot the performance ratio of a subset of the remaining heuristics over that of the L.P. with adaptive period: in Figure 6, we use a 5 processor platform. For

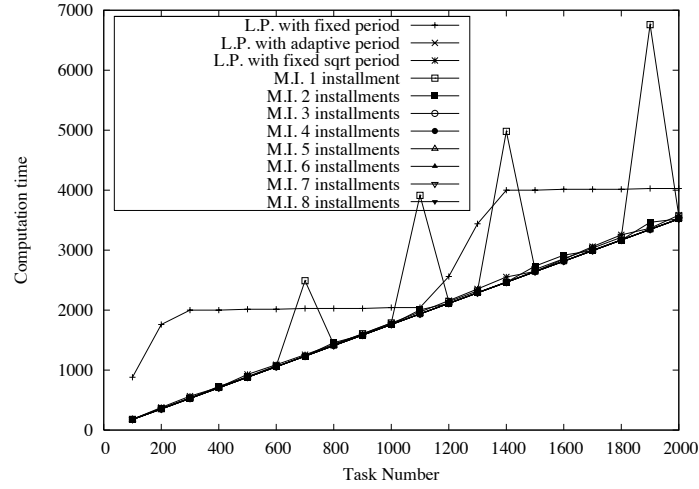


Figure 5: Simulation time for an homogeneous platform with 5 processors.

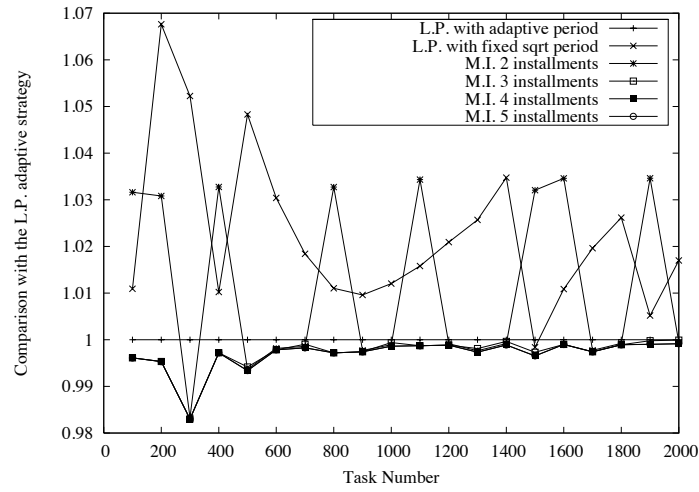


Figure 6: Comparison with the adaptive heuristic for an homogeneous platform with 5 processors.

the sake of clarity, only M.I.x with $x = 2$ to $x = 5$ are reported in the figure; no improvement is obtained with larger values of x . The following observations can be made:

- the adaptive strategy is always very close to the best heuristic;
- the number of rounds of the Multi Installment algorithm must be at least 3, and no real improvement is to be expected when increasing that number;
- the linear programming with fixed square-root period is not very regular but stays within 5% of the optimal heuristic, as soon as the number of tasks (work units) grows sufficiently.

The behavior is similar when increasing the size of the platform. The adaptive strategy may need a larger task set to achieve the same efficiency. The number of installments needed to obtain good results with the Multi Installment algorithm also becomes higher. As an example, Figure 7 depicts the results for a 20-processor platform.

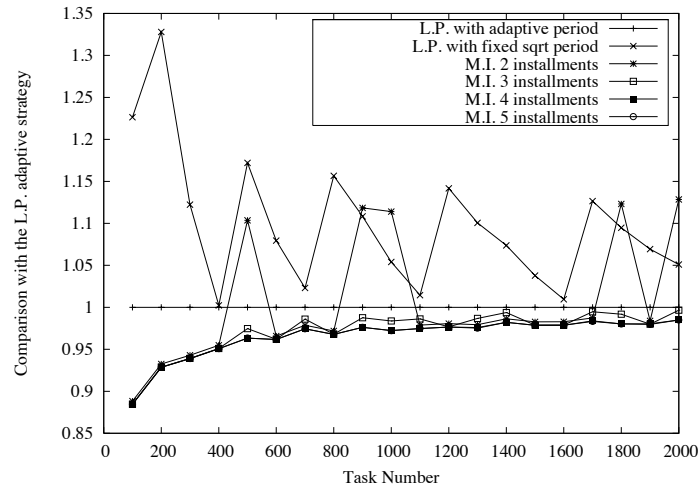


Figure 7: Comparison with the adaptive approach for an homogeneous platform with 20 processors.

6.2 Heterogeneous platforms, no latency

Experiments have been conducted with 2000 simulated platforms made up of machines randomly chosen in the following processor set : PPro 200MHz (22.151 Mflops), PII 450MHz (48.492 Mflops), PII 350MHz (34.333 Mflops), and PIII 1GHz (114.444 Mflops). The network used to interconnect the slaves to the master could be Ethernet either 10 Mbits/secs or 100 Mbits/sec (we measured 4.70, 32.10 or 30.25 Mbits/sec of effective bandwidth). The number of workload units ranges from 100 to 2000; with a step of 100. As before, one task unit amounts to 1 GFlops of computation and 2 Mbits of data. In other words, the workers communicate during between 0.06 and 0.4 seconds to be able to compute during between 9 and 90 seconds. Again, the communication-to-computation ratio is quite low, which makes it easier to obtain good performances.

We have run one simulation per platform and per task set, which amounts to 40000 experiments per heuristic; the figures below correspond to averaged values. Figure 8 and its

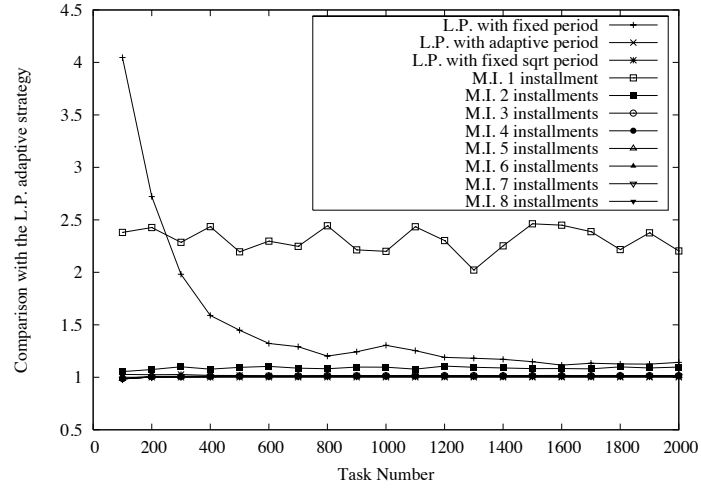


Figure 8: Comparison with the adaptive approach for heterogeneous platforms with 5 processors.

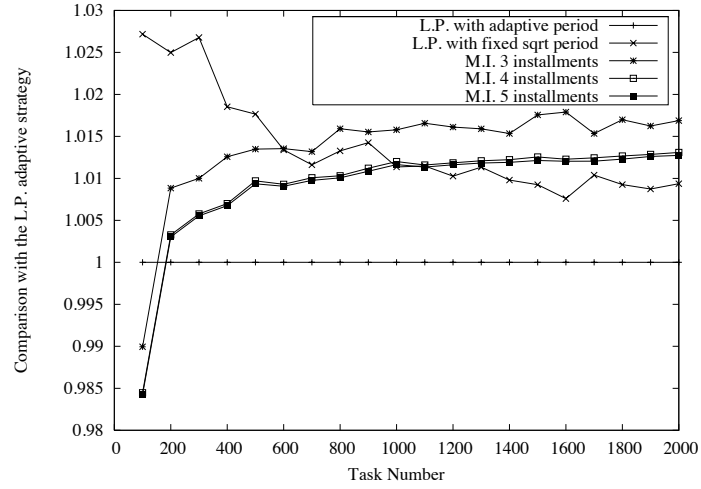


Figure 9: Comparison with the adaptive approach for heterogeneous platforms with 5 processors (zoom).

zoomed counterpart Figure 9 depict the comparison of the heuristics with the adaptive linear programming approach on 5-processor platforms. The following observations can be made:

- linear programming with fixed period, and one installment strategies lead to very poor performances;
- the optimal number of rounds of the Multi Installment algorithm is close to 4 (for the sake of clarity, only M.I.x with $x = 2$ to $x = 5$ installments are depicted on Figure 9 because increasing the number of rounds does not improve the performances);
- the linear programming with fixed square-root period slightly outperforms the multi installment algorithm;
- the adaptive approach leads to the best performances but with a small improvement over the other good heuristics (1% in average).

When increasing the size of the platform, the optimal number of rounds of the multi-installment gets larger. Linear programming strategies show good performances only when the task set is large enough; the adaptive method remains better than other linear programming approaches. As an example, Figure 10 depicts the results for 20-processor platforms.

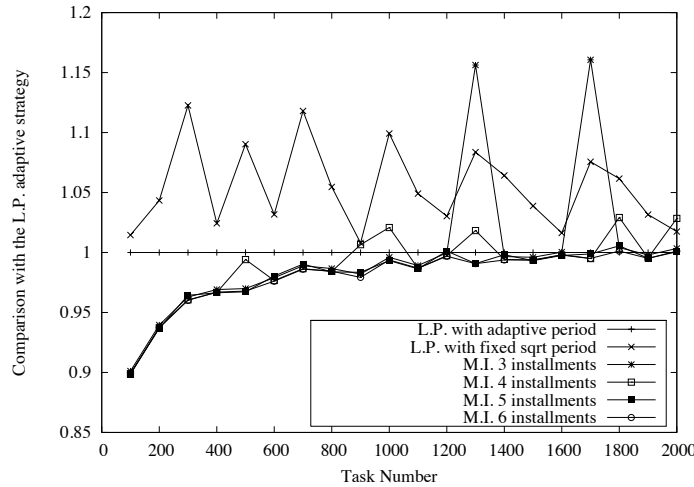


Figure 10: Comparison with the adaptive approach for heterogeneous platforms with 20 processors.

6.3 Heterogeneous platforms, with latency

6.3.1 With a low communication-to-computation ratio

Experiments have been conducted with the same platform set and the same workloads as in Section 6.2. The only difference resides in the fact that we add 2 Mbits of data to each message, so as to model the latencies: in other words, we set $g_i = 2 \cdot 10^6 \cdot G_i$ for every worker P_i .

Figure 11 and its zoomed counterpart Figure 12 depict the comparison of the heuristics with the adaptive linear programming approach on 5-processor platforms. Multi Installments

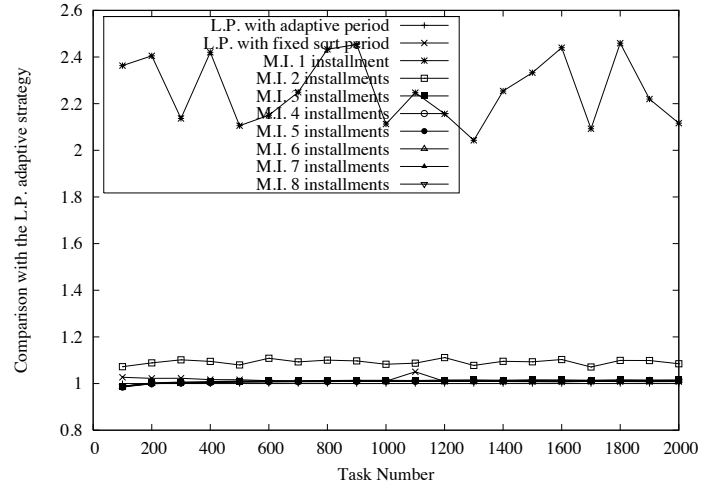


Figure 11: Comparison with the adaptive approach for heterogeneous platforms with 5 processors, with latencies.

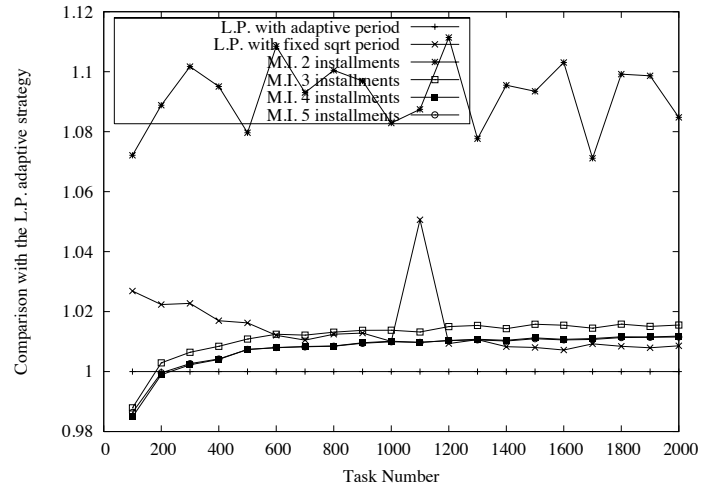


Figure 12: Comparison with the adaptive approach for heterogeneous platforms with 5 processors, with latencies (zoom).

with a small number of rounds still lead to poor performances. The optimal number of rounds for the Multi Installment algorithm is equal to 4. The best strategy is the adaptive approach, which still leads to a small improvement of 1%. When the size of the platform gets larger (see Figure 13), the optimal number of rounds for the Multi Installment algorithm is equal to 5 and the adaptive strategy is at most 3% far away from the best other solutions (when the number of task is larger than 200).

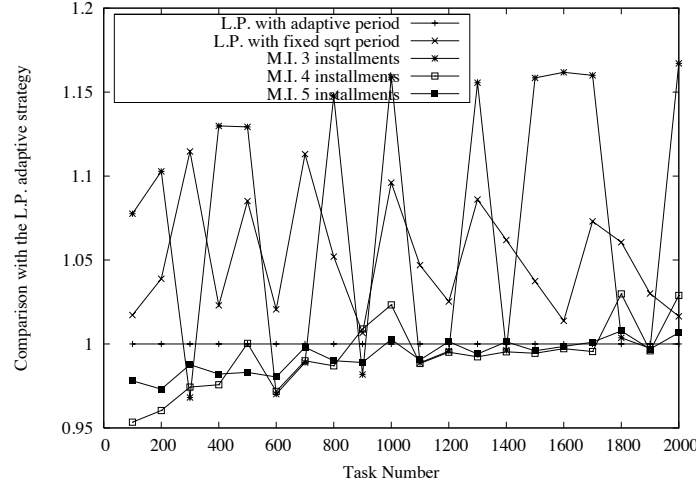


Figure 13: Comparison with the adaptive approach for heterogeneous platforms with 20 processors, with latencies.

6.3.2 With a high communication-to-computation ratio

Experiments have been conducted with, again, the same platform set as in Section 6.2, the same workloads and the same latencies. But we change the communication-to-computation ratio: one task unit now amounts to 50 MFlops of computation and, still, 2 Mbits of data: the communication-to-computation ratio has roughly been multiplied by 20.

Figures 14 and 15 depict the comparison of the heuristics with the adaptive linear programming approach on 5-processor and 20-processor platforms. Linear programming approaches clearly outperform any Multi Installment heuristic. This is due to many reasons:

- The linear equations given in [6] cannot take latencies into account. Since the communication-to-computation ratio is getting higher, considering them becomes crucial;
- No resource selection is done in [6]. When the size of the platform gets larger, the network becomes a bottleneck and we must decide which computing resources to use. This choice is automatically performed when solving the linear inequalities.

6.4 Summary

As a general conclusion, we see that L.P. strategies, either with fixed square-root period, or with the adaptive strategy to compute the next period, are to be recommended. In most cases, they are very close to the best multi-installment solution (and determining the optimal number

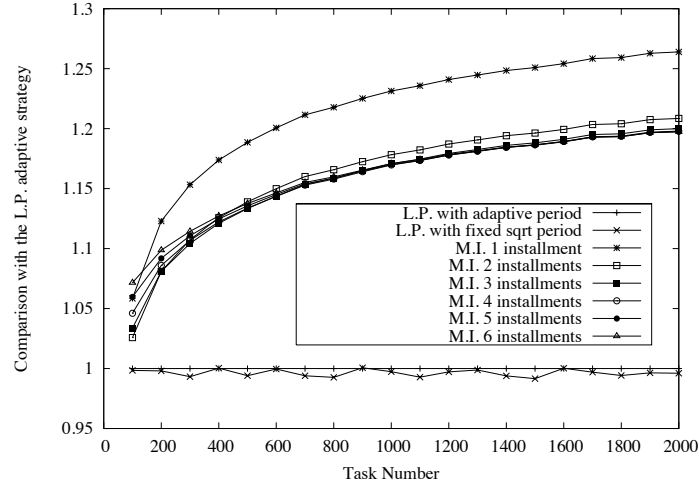


Figure 14: Comparison with the adaptive approach for heterogeneous platforms with 5 processors, with latencies.

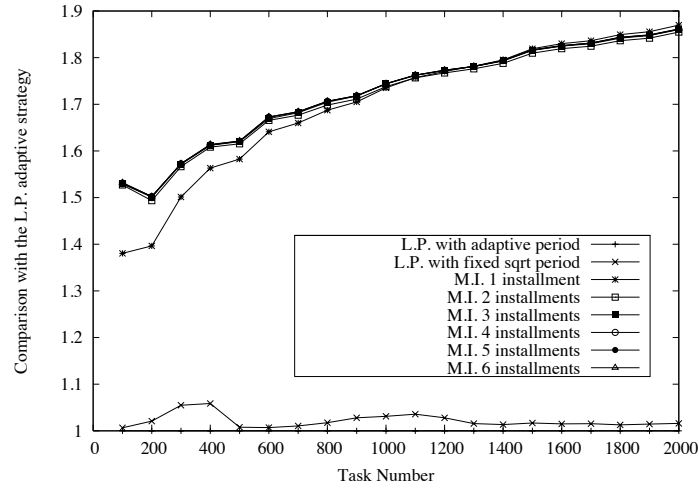


Figure 15: Comparison with the adaptive approach for heterogeneous platforms with 20 processors, with latencies.

of rounds for this class of heuristics is a non-trivial problem). When the communication-to-computation ratio gets higher, both L.P. strategies are much better than the other heuristics.

If we were to make a final choice, we would recommend the adaptive strategy, because it would be the most robust to changes in computation speeds or network bandwidths.

7 Conclusion

On the theoretical side, the main result of this paper is the proof of the asymptotic optimality of our multi-round algorithm. This is the first quantitative result ever assessed for a multi-round algorithm. But (maybe more importantly), our algorithm exhibits a lot of interesting features that make it a candidate of choice in a wide variety of situations:

- The best selection of the resources to be used among all available machines is automatically conducted through the linear program. Even better, resources are sorted according to the G_i and greedily selected until the sum of the ratios $\frac{G_i}{G_i + w_i}$ (without overlap) or $\frac{G_i}{w_i}$ (with overlap) exceeds 1. Previous approaches had to resort to un-guaranteed experimental heuristics.
- The best number of rounds is easily determined as a function of the task number, so there is no need to test several solutions with different round numbers, and then to select the best one.
- Because it is periodic, the algorithm is simpler to implement than other schemes that grow the chunk size repeatedly.
- For the same reason, our algorithm is more robust: the decisions taken for each round (how many work units should be sent to each worker) can be questioned before each round, thus allowing a dynamic approach to cope with, and respond to variations in computation speeds or network bandwidths. Such changes are very likely to occur, especially when the overall processing time is large. Other algorithms that rely on very long rounds in the end cannot rapidly adapt to speed or bandwidth changes.

Extensive simulations have shown that our multi-round algorithm does perform very well in practice, and significantly outperforms other heuristics when the communication-to-computation ratio of the application is not too low on the target platform. This opens up a larger range of applications for the divisible workload paradigm.

References

- [1] D. Altilar and Y. Paker. An optimal scheduling algorithm for parallel video processing. In *IEEE Int. Conference on Multimedia Computing and Systems*. IEEE Computer Society Press, 1998.
- [2] D. Altilar and Y. Paker. Optimal scheduling algorithms for communication constrained parallel processing. In *Euro-Par 2002*, LNCS 2400, pages 197–206. Springer Verlag, 2002.
- [3] G. Ausiello, P. Crescenzi, G. Gambosi, V. Kann, A. Marchetti-Spaccamela, and M. Prota. *Complexity and Approximation*. Springer, Berlin, Germany, 1999.

- [4] S. Bataineh, T.Y. Hsiung, and T.G. Robertazzi. Closed form solutions for bus and tree networks of processors load sharing a divisible job. *IEEE Transactions on computers*, 43(10):1184–1196, October 1994.
- [5] O. Beaumont, L. Carter, J. Ferrante, A. Legrand, and Y. Robert. Bandwidth-centric allocation of independent tasks on heterogeneous platforms. In *International Parallel and Distributed Processing Symposium IPDPS'2002*. IEEE Computer Society Press, 2002. Extended version available as LIP Research Report 2001-25.
- [6] V. Bharadwaj, D. Ghose, V. Mani, and T.G. Robertazzi. *Scheduling Divisible Loads in Parallel and Distributed Systems*. IEEE Computer Society Press, 1996.
- [7] J. Blazewicz, M. Drozdowski, and M. Markiewicz. Divisible task scheduling - concept and verification. *Parallel Computing*, 25:87–98, 1999.
- [8] H. Casanova. Simgrid: A toolkit for the simulation of application scheduling. In *Proceedings of the IEEE Symposium on Cluster Computing and the Grid (CCGrid'01)*. IEEE Computer Society, May 2001.
- [9] H. Casanova and F. Berman. Grid'2002. In F. Berman, G. Fox, and T. Hey, editors, *Parameter sweeps on the grid with APST*. Wiley, 2002.
- [10] B. W. Char, K. O. Geddes, G. H. Gonnet, M. B. Monagan, and S. M. Watt. *Maple Reference Manual*, 1988.
- [11] S. Charcraonon, T.G. Robertazzi, and S. Luryi. Optimizing computing costs using divisible load analysis. *IEEE Transactions on computers*, 49(9):987–991, September 2000.
- [12] P. Chrétienne, E. G. Coffman Jr., J. K. Lenstra, and Z. Liu, editors. *Scheduling Theory and its Applications*. John Wiley and Sons, 1995.
- [13] D. E. Culler and J. P. Singh. *Parallel Computer Architecture: A Hardware/Software Approach*. Morgan Kaufmann, San Francisco, CA, 1999.
- [14] M. Drozdowski. *Selected problems of scheduling tasks in multiprocessor computing systems*. PhD thesis, Instytut Informatyki Politechnika Poznańska, Poznań, 1997.
- [15] H. El-Rewini, T. G. Lewis, and H. H. Ali. *Task scheduling in parallel and distributed systems*. Prentice Hall, 1994.
- [16] M. R. Garey and D. S. Johnson. *Computers and Intractability, a Guide to the Theory of NP-Completeness*. W. H. Freeman and Company, 1991.
- [17] T. Hagerup. Allocating independent tasks to parallel processors: an experimental study. *J. Parallel and Distributed Computing*, 47(2):185–197, 1997.
- [18] S. Flynn Hummel, E. Schonberg, and L.E. Flynn. Factoring: a method for scheduling parallel loops. *Communications of the ACM*, 35(8):90–101, 1992.
- [19] C. Lee and M. Hamdi. Parallel image processing applications on a network of workstations. *Parallel Computing*, 21:137–160, 1995.

- [20] J. Lerouge and A. Legrand. Towards realistic scheduling simulation of distributed applications. Technical Report 2002-28, LIP, ENS Lyon, France, July 2002.
- [21] A. L. Rosenberg. Sharing partitionable workloads in heterogeneous NOWs: greedier is not better. In D. S. Katz, T. Sterling, M. Baker, L. Bergman, M. Paprzycki, and R. Buyya, editors, *Cluster Computing 2001*, pages 124–131. IEEE Computer Society Press, 2001.
- [22] J. Sohn, T.G. Robertazzi, and S. Luryi. Optimizing computing costs using divisible load analysis. *IEEE Transactions on parallel and distributed systems*, 9(3):225–234, March 1998.
- [23] R.Y. Wang, A. Krishnamurthy, R.P. Martin, T.E. Anderson, and D.E. Culler. Modeling communication pipeline latency. In *Measurement and Modeling of Computer Systems (SIGMETRICS'98)*, pages 22–32. ACM Press, 1998.
- [24] Y. Yang and H. Casanova. Multi-round algorithm for scheduling divisible workload applications: analysis and experimental evaluation. Technical Report CS2002-0721, Dept. of Computer Science and Engineering, University of California, San Diego, 2002.